

当代经济学系列丛书

Contemporary Economics Series

主编 陈昕

CAMBRIDGE

政治博弈论

当代经济学
教学参考书系

[美] 诺兰·麦卡蒂 亚当·梅罗威茨 著

孙经纬 高晓晖 译



格致出版

格致出版社
上海三联书店
上海人民出版社

VISIT...

LANZAROTE
Caliente.COM



政治博弈论

终于有了一本专为政治科学家学习博弈论而撰写的研究生教材，它既有深度又适合学习。这本书用政治科学中的例子讲授博弈论，将是未来许多年中的标准教材。

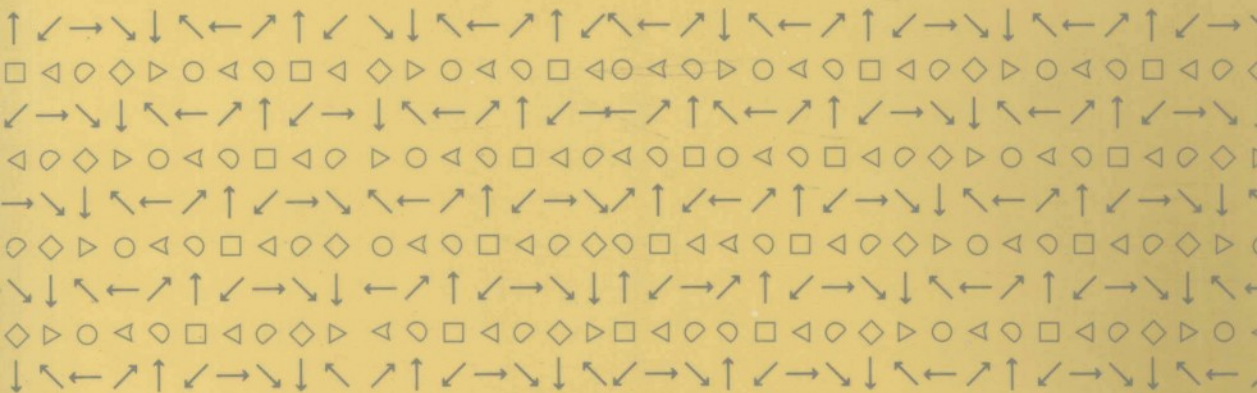
——Robert Powell，加州大学伯克利分校

麦卡蒂和梅罗威茨为政治科学家撰写了一本重要的博弈论教材。它精雕细琢，在技术上十分细致，同时又很适合学习。在著名大学的政治科学系中，它将成为研究生的标准教材。

——Ken Shepsle，哈佛大学

《政治博弈论》对现代博弈论和它在政治科学中的应用作了最全面的介绍。对这一领域中的学生和教师，它应该是本重要参考书。

——Michael Ting，哥伦比亚大学



CAMBRIDGE
UNIVERSITY PRESS
www.cambridge.org



ISBN 978-7-5432-1668-6



9 787543 216686 >

定价：45.00元

易文网：www.ewen.cc

格致网：www.hibooks.cn

图书在版编目(CIP)数据

政治博弈论/(美)麦卡蒂,(美)梅罗威茨著;孙
经纬,高晓晖译. —上海:格致出版社:上海人民出版社,
2009

(当代经济学系列丛书. 当代经济学教学参考书系)

ISBN 978-7-5432-1668-6

I. 政… II. ①麦…②梅…③孙…④高… III. 政治理论
IV. D0

中国版本图书馆 CIP 数据核字(2009)第 167401 号

责任编辑 谷 雨 钱 敏
装帧设计 敬人设计工作室
吕敬人

政治博弈论

[美]诺兰·麦卡蒂 亚当·梅罗威茨 著
孙经纬 高晓晖 译

格致出版社·上海三联书店·上海人民出版社
(200001 上海福建中路 193 号 24 层 www.ewen.cc)



编辑部热线 021-63914988
市场部热线 021-63914081
www.hibooks.cn

世纪出版集团发行中心发行
上海图字印刷有限公司印刷

2009 年 11 月第 1 版
2009 年 11 月第 1 次印刷
开本:787×1092 1/16
印张:20.5 插页:5 字数:414,000

ISBN 978-7-5432-1668-6/F·217

定价:45.00 元

Political Game Theory: An Introduction (978-0-521-84107-8) by Nolan McCarty and Adam Meirowitz first published by Cambridge University Press 2007.

All rights reserved.

This simplified Chinese edition for the People's Republic of China is published by arrangement with the Press Syndicate of the University of Cambridge, NY, United States of America.

© Cambridge University Press & Shanghai People's Publishing House 2009

This book is in copyright. No reproduction of any part may take place without the written permission of Cambridge University Press or Shanghai People's Publishing House.

This edition is for sale in the mainland of China only, excluding Hong Kong SAR, Macao SAR and Taiwan, and may not be bought for export therefrom.

此版本仅限中华人民共和国境内销售, 不包括中国香港、澳门特别行政区及中国台湾地区。不得出口。

上海市版权局著作权合同登记号 图字 09-2007-828



出版前言

为了全面地、系统地反映当代经济学的全貌及其进程,总结与挖掘当代经济学已有的和潜在的成果,展示当代经济学新的发展方向,我们决定出版“当代经济学系列丛书”。

“当代经济学系列丛书”是大型的、高层次的、综合性的经济学术理论丛书。它包括三个子系列:(1)当代经济学文库;(2)当代经济学译库;(3)当代经济学教学参考书系。该丛书在学科领域方面,不仅着眼于各传统经济学学科的新成果,更注重经济前沿学科、边缘学科和综合学科的新成就;在选题的采择上,广泛联系海内外学者,努力开掘学术功力深厚、思想新颖独到、作品水平拔尖的“高、新、尖”著作。“文库”力求达到中国经济学界当前的最高水平;“译库”翻译当代经济学的名人名著;“教学参考书系”则主要出版国外著名高等院校的通用教材。

本丛书致力于推动中国经济学的现代化和国际标准化,力图在一个不太长的时期内,从研究范围、研究内容、研究方法、分析技术等方面逐步完成中国经济学从传统向现代的转轨。我们渴望经济学家们支持我们的追求,向这套丛书提供高质量的标准经济学著作,进而为提高中国经济学的水平,使之立足于世界经济之林而共同努力。

我们和经济学家一起瞻望着中国经济学的未来。

致 谢

写作本书的初衷,是本书的一位作者完全无法在黑板上(或者在任何平面上)正确地板书。为使学
生免于在听课过程中忍受这种折磨,我们把讲义制作成电子版。而且,如果没有拼写检查器的帮助,我们还
没办法正确拼写英文单词,所以,电子版讲义弥补了
我们在这个方面的欠缺。^①对耗费许多个夜晚输入
电脑中的这份博弈论讲义,我们觉得不应该束之高
阁。所以,就把它变为了这本书,尽管这又使我们更
频繁地打字到深夜。我们希望这份辛苦没有白费。

我们感谢哥伦比亚大学和普林斯顿大学的学生,他们使用了讲义和手稿的各个早期版本。正是他们的迷惑不解和办公室中的答疑,帮助我们学习如何向政治学系学生讲授博弈论。在政治博弈论教材应该是怎样的一本书的问题上,我们受益于与许多人的交流,包括 Chris Achen、Scott Ashworth、Larry Bartels、Cathy Haffer、Keith Krehbiel、David Lewis、Kris Ramsay 和 Thomas Romer。我们感谢 John Londregan 和 Mark Fey,他们指出了此书早期的各个版本中的错误。我们还要感谢的是教授我们政治博弈论的老师:David Austen-Smith、Jeffrey Banks、David Baron、Bruce Bueno de Mesquito、Thomas Romer 和 Howard Rosenthal。

诺兰·麦卡蒂
亚当·梅罗威茨

① 但是,我们两人的拼写错误又是两种不同的风格。对一个单词,McCarty 完全在随机拼写,Meirowitz 则始终以相同的方式错误地拼写。

目 录

001	出版前言
001	致谢
001	1 导论
002	1.1 本书的结构
004	2 选择理论
005	2.1 有限行动集和有限结果集
007	2.2 连续选择空间
013	2.3 效用理论
014	2.4 连续选择空间上的效用表示
015	2.5 空间偏好
017	2.6 习题
019	3 不确定性下的选择
019	3.1 有限行动集和有限结果集
027	3.2 风险偏好
033	3.3 学习
037	3.4 对期望效用理论的批判
042	3.5 时间偏好
046	3.6 习题
049	4 社会选择理论
049	4.1 公开招聘
050	4.2 偏好加总规则

新华书店
PDG

056	4.3 集体选择
061	4.4 操控选择函数
063	4.5 习题
065	5 标准式博弈
067	5.1 标准式博弈
070	5.2 标准式博弈的解
075	5.3 应用: Hotelling 政治竞争模型
080	5.4 Nash 均衡的存在性
085	5.5 占优和混合策略
086	5.6 计算 Nash 均衡
088	5.7 应用: 利益集团献金
089	5.8 应用: 国际外部性
091	5.9 利用约束条件下的优化方法, 计算均衡
092	5.10 证明 Nash 均衡的存在性
095	5.11 比较静态
103	5.12 精炼 Nash 均衡
104	5.13 应用: 私人提供公共品
108	5.14 习题
112	6 标准式贝叶斯博弈
113	6.1 正式定义
115	6.2 应用: 贸易保护
116	6.3 应用: 陪审团投票
119	6.4 应用: 在信号集为连续统时陪审团的投票行为
120	6.5 应用: 公共品和不完全信息
123	6.6 应用: 候选人偏好上的不确定性
124	6.7 应用: 竞选、竞赛和拍卖
126	6.8 贝叶斯 Nash 均衡的存在性
127	6.9 习题
128	7 扩展式博弈
131	7.1 反向归纳法
132	7.2 完全非完美信息动态博弈
138	7.3 单偏离原则
138	7.4 子博弈精炼和精炼均衡
139	7.5 应用: 议程控制

143	7.6 应用:力量结构变化模型
145	7.7 应用:向民主制度的转轨
148	7.8 应用:政党联盟的形成模型
151	7.9 习题
153	8 不完全信息动态博弈
156	8.1 精炼贝叶斯均衡
160	8.2 信号揭示博弈
164	8.3 应用:选举中的阻止进入行为
170	8.4 应用:信息和立法组织
174	8.5 应用:信息沟通目的的游说活动
177	8.6 对精炼贝叶斯均衡的精炼
185	8.7 习题
189	9 重复博弈
190	9.1 重复的囚徒两难博弈
191	9.2 触发均衡
192	9.3 以牙还牙策略
194	9.4 介于触发策略和以牙还牙策略之间的惩罚策略
196	9.5 无名氏定理
197	9.6 应用:团体间合作
202	9.7 应用:贸易战
205	9.8 习题
207	10 讨价还价理论
207	10.1 Nash 讨价还价解
211	10.2 非合作讨价还价
216	10.3 封闭规则下多数通过的讨价还价
221	10.4 开放规则下的 Baron-Ferejohn 模型
223	10.5 不完全信息下的讨价还价
224	10.6 应用:包含否决行为的讨价还价
233	10.7 应用:危机谈判
241	10.8 习题
242	11 机制设计和代理理论
243	11.1 例子
244	11.2 机制设计问题

新华书店
PDG

246	11.3 应用: 民调
248	11.4 拍卖理论
252	11.5 应用: 竞选和全支付拍卖
255	11.6 激励一致性和个人理性
257	11.7 约束条件下的机制设计
271	11.8 机制设计和信号揭示博弈
275	11.9 习题
277	12 数学附录
277	12.1 数学命题和证明
279	12.2 集合和函数
282	12.3 实数
284	12.4 点和集合
285	12.5 函数的连续性
287	12.6 对应
288	12.7 微积分
303	12.8 概率论
313	参考文献

导论

在很短的时间里,博弈论成为研究政治问题的最有力的分析工具。从最早被用于研究选举和立法行为开始,博弈论模型已经广泛应用于诸如国际安全、种族合作和民主化等各种领域的研究中。实际上,政治科学所有领域都受益于博弈论模型的重大贡献。在《美国政治科学评论》(*American Political Science Review*)、《美国政治科学期刊》(*American Journal of Political Science*)和《国际组织》(*International Organization*)等刊物中,几乎每一期都有至少一篇论文对政治现象提出新的博弈模型分析,或者有至少一篇论文对已有模型进行经验实证检验。

但是,博弈论在政治科学中的应用并不像在经济学中那样发展迅速。造成这种发展上的不平衡的原因之一,是学习博弈论的大部分政治科学家只能依靠经济学家撰写的和为经济学家撰写的教科书。有许多杰出的经济学博弈论教材,但是,所探讨的问题常常并不切合政治科学家的需要。首先,可能也是最重要的,这些书中的博弈论应用和所探讨的主题,一般是经济学家们感兴趣的内容。例如,新入行的政治科学家们并不总能够直接看到双寡头理论或拍卖理论对分析政治现象具有的价值。其次,对政治博弈论学家们来说,一些重要的问题,如投票理论等,在经济学教材中极少涉及。许多经济学教材都假定读者对经典价格理论有一定的了解。结果,对政治科学家们来说,进入壁垒不仅包含数学,而且包括需求曲线和边际替代率等知识。

当然了,我们有几本由政治科学家撰写的和为政治科学家撰写的教材,例如 Ordeshook(1986)和 Morrow(1994)。但是,我们认为,无论就应用而言,还是就现代政治科学的需要而言,这些教材都有些过时。Ordeshook 的教材对社会选择和空间理论的探讨,依然独领风骚,但是,它是在非合作理论成为政治博弈论的主导范式之前撰写和出版的。Morrow 对非合作博弈论的阐述通俗易懂,其分析深度却无法满足当代学者的需要。而且,它自出版之今已有 10 年时间——在这 10 年中,出现了数以百计的重要论文和著作使用博弈论工具展开分析。Austen-Smith 和 Banks 的著作(1999, 2005)部分地满足了这方面的需要。他们的第一本书,《实证政治理论 I》(*Positive Political Theory I*)对社会选择理论进行了透彻研究,而在本书中,我们仅用一章的篇幅介绍这一领域。他们的第二本书,《实证政治理

论Ⅱ》(Positive Political Theory Ⅱ),虽探讨策略和制度,却假定读者了解博弈论,而这是政治科学领域中一年级学生所不具备的知识。而且,他们的这本书是按照研究专题组织的,而不是按照博弈论体系展开的。

所以,在撰写本书时,我们确定了几个目标。首先,我们要撰写的是一本政治博弈论教材,而不是抽象的博弈论或经济博弈论的教材。因此,我们在本书中所探讨的博弈论的应用,都是政治科学家们感兴趣的应用;所选择的内容,也是政治分析所独有的。其次,在为政治科学家们撰写这样一本书时,我们想适应青年政治科学家在背景和需要上的多样性。我们知道,在政治科学专业中,大部分博士生在就读时,只有有限的数学基础和模型分析基础。但是,我们认为,因此而忽视当代政治模型所依托的数学上的严谨性和关键的理论概念,对他们不是件好事。对那些基础欠缺的学生,我们在本书中有专门而详细的数学附录,提供了从集合论和分析到基本优化和概率论的一些必需的工具。当然了,有一些攻读政治科学专业的研究生,有较强的数学和经济学基础。我们希望这本书对他们同样有用。因此,对一些较难和较复杂的概念,我们也做了深入阐述。许多高深的章节(用“*”或“**”标记)对我们所探讨的模型的分析结构和数学结构,提供了详细阐述。对于尚没有为这些技术性更强的部分做好基础准备的学生,在第一次接触本书时,完全可以跳过这些章节。

1.1 本书的结构

本书在内容的组织上,偏离了标准做法,因为其中包含的许多内容或者直接与政治科学有关,或者是针对政治科学专业的学生薄弱的知识领域,提供补充。

第2章对确定性条件下经典选择理论提供了完整阐述。我们介绍了偏好和效用理论的基本思想。我们证明了一些重要结论;其中一些证明十分简单,还有一些证明则用到了比较高深的数学知识,对后者,我们用“*”做了标记。本章的重心是介绍理性选择理论的内涵和所使用的语言。我们专门用了一节的篇幅介绍空间偏好或者说欧几里得偏好。在投票理论和在投票理论于选举政治和立法政治的应用中,这类偏好起着核心作用。

在第3章中,我们介绍了博弈理论家们是如何分析不确定性条件下的选择的。本章的核心是标准的 von Neumann-Morgenstern 期望效用模型,同时介绍了对这一理论的一些批判。在探讨了风险偏好之后,我们还讨论了在行为主体有空间偏好时,风险因素所具有的特殊影响。

第4章简短介绍了社会选择理论。这一章的目的不是取代 Peter Ordeshook (1986)以及 David Austen-Smith 和 Jeff Banks(1999)等社会选择教材,而主要是为已经成为数理政治科学不可分割部分的概念和思想,提供参考资料。这些思想和概念包括 Arrow 不可能性定理、多数通过规则的核为空集/非空集以及单峰偏好的作用等。本章还讨论了关于社会决策中的策略性行为的 Gibbard-Satther-

waite 定理。

第5章开始讨论当代数理政治理论的核心分析工具：非合作博弈论。我们介绍了完全信息下的标准式博弈，展示了其最基本的解概念：占优和 Nash 均衡。这部分的理论阐述相当标准，同时，还介绍了一些重要的政治学应用：介绍了关于竞选的 Downsian 模型以及 Donald Wittman 和 Randy Calvert 对这一模型的发展；之后，以 Thomas Palfrey 和 Howard Rosenthal 的论文为基础，介绍了几个公共品生产中的私人出资模型。

第6章发展了标准式博弈，分析了参与者并不确切了解各个策略组合所产生的收益的情况。在介绍了这类博弈的解概念即贝叶斯 Nash 均衡后，探讨了第5章中各模型的不完全信息版本。这种比较有助于理解不确定性的策略意义。

第7章探讨完全信息下的多阶段动态博弈，提出了子博弈精炼概念。在这一章中，我们还将介绍动态博弈在立法政治、民主转轨、联盟的形成和国际危机谈判中的应用。

在第8章所介绍的动态博弈中，某些参与者在各种策略组合所实现的收益上只拥有不完全信息。在探讨如何求解这种模型后，我们还将介绍这种博弈在立法政治、竞选融资和国际谈判中的应用。我们还用很大篇幅介绍信号揭示博弈，这类博弈在政治科学中有着越来越重要的应用。

第9章则介绍了重复博弈理论和它在政治科学中的应用，所探讨的重点是时间贴现的作用和无名氏定理的结构。这类模型的应用为团体间合作和贸易战。

第10章探讨讨价还价理论的应用。我们首先介绍了标准的 Nash 讨价还价模型和 Rubinstein 的讨价还价模型，然后介绍了 Baron 和 Ferejohn 提出的多数通过的投票规则下的讨价还价博弈。最后介绍了不完全信息下讨价还价理论的几个应用。

第11章阐述的是分析制度所使用的机制设计理论。我们探讨的中心问题是如何通过对博弈形式的选择，诱生出满足某些目标的均衡行为。在介绍了揭示原理和激励一致性条件后，我们介绍了一些近期文献，包括机制设计在选举政治和组织设计中的应用。在第8章的基础上，我们还找出了信号揭示模型和机制设计理论之间的一些联系。

为了使本书自成体系，第12章介绍了本书中用到的所有数学知识。这些知识是本书中重要的理论结论和应用分析不可缺少的部分，所涉及的数学领域包括集合论、实分析、线性代数、微积分、优化和概率论。实际上，这一章既可以用于复习，也可以自学。想在政治博弈论前沿领域中作一番研究工作的学生，如果想娴熟地掌握这一章中的数学概念，还得学习更多的数学课程。

▶ 2

选择理论

在物质资源和对其他人行为的预期所施加的约束下,人们理性地追求自己的目标。这一假定是政治博弈论的基础。但是,在这一假定上人们又常常争执不休。事实上,社会科学中的最活跃的讨论之一,便是理性和目的性在预测人们的行为中所起的作用。这里不讨论经济人和社会人两个概念上的争论,而直接介绍理性选择的基本模型。

就我们的目的而言,我们只需把理性定义在以下两个特征的基础之上:

(1) 面对着任意两个选项 x 和 y , 每个人都能够确定自己是不偏好 x 、不偏好 y 还是对其中任何一个选项都不偏好。在偏好满足这一特征时,我们就说偏好是完全的。

(2) 面对着任意三个选项 x 、 y 和 z , 如果在 x 和 y 之间某个人不偏好 y 且在 y 和 z 之间不偏好 z , 则在 x 和 z 之间, 他肯定不偏好 z 。具有这一特征的偏好, 是传递的。

我们把理性行为定义为与完全的和传递的偏好相一致的行为。我们有时也称这种行为是弱理性的(thinly rational), 因为上述两个特征没有谈到人的欲望。与弱理性相对应的是强理性(thick rationality)。在使用强理性概念时, 分析者会明确人们的目标, 如财富、地位或名声等。弱理性概念支持各种各样的目标。理论上说, 各种因素——意识形态、价值观甚至宗教等——都可以是具有弱理性的人追求的目标。只要这些信仰体系, 在个人的选择集上或者在社会的选择集上, 产生完全的和传递的偏好关系, 我们就可以用选择理论分析人们的行为。

不对实际目标施加明确的假定的做法, 尽管有吸引力, 但是, 我们常常必须赋予偏好更强的假定——例如利益集团最大化其成员的财富, 政治家最大化其再次当选的概率等。在本书讨论的模型中, 我们都对人们的偏好施加这类假定。实际上, 理性模型对分析各种类型的活跃分子的行为——例如, 出于信念和非物质财富因素, 这些人想最小化环境恶化程度或堕胎数量等——亦十分有用。

本章介绍确定性条件下的选择理论。在确定性条件下, 决策者拥有关于其行动集的充分信息, 因此能完全预测其每个行动导致的结果。下一章则分析不确定性条件下的选择——决策者缺少信息而只能在所导致的结果并不确定的行动之间做出选择。

2.1 有限行动集和有限结果集

首先看一个简单的选择问题：某个人要从数量有限的行动中，选择一项行动。我们将可供他从中做出选择的行动表示为集合 $A = \{a_1, \dots, a_k\}$ 。在国际危机中，国家领导人面对着的行动集可能为 $A = \{\text{出兵, 谈判, 不作为}\}$ 。在美国的总统选举中，选民的行动集可能为 $A = \{\text{投票给民主党, 投票给共和党, 弃权}\}$ 。

如前所述，我们假定决策者拥有完全信息——他们拥有充分信息，因而能准确预测到每个行动导致的结果。为正式阐述这一假定，我们把结果集定义为 $X = \{x_1, \dots, x_n\}$ 。在国际危机例子中， $X = \{\text{大胜却损兵折将, 小胜, 现状}\}$ 。确定性意味着每个行动 $a \in A$ 直接对应着一个且唯一一个 $x \in X$ 。就是说，确定性意味着存在函数 $x: A \rightarrow X$ ，它把每个行动映射到某个结果上。为便于分析，我们假定集合 X 中的所有结果都是可行的，就是说，给定每个结果，存在着至少一个行动导致这一结果。换句话说， x_i 是可行的，如果存在 $a \in A$ 使得 $x(a) = x_i$ 。在确定性和可行性假设下，决策者在行动上的偏好和在结果上的偏好，没有什么区别。所以，本章只讨论决策者在结果上的偏好。在第3章中，不确定性假设或者说不完全信息假设，使得行动和结果之间的区别变得重要。

要预测决策者的行为，我们还需要更正式的偏好概念。我们定义弱偏好为两元关系 R ：如果 $x_i R x_j$ ，则在结果 x_i 和 x_j 之间，决策者不偏好 x_j 。换言之，如果 $x_i R x_j$ ，则在 x_i 和 x_j 之间，决策者“弱”偏好 x_i 。^①这里的 R ，类似实数集上的两元关系“ $\geq$ ”（大于或等于）。

利用弱偏好关系 R ，我们定义另外两种重要的两元关系：严格偏好关系和无差异关系。

定义 2.1 对任意的 $x, y \in X$ ， xPy (x 严格优于 y) 当且仅当 xRy 且无 yRx ； xIy (x 与 y 无差异) 当且仅当 xRy 且 yRx 。

在这里， P 表示严格偏好， I 表示无差异。如果把 R 类比为实数集上的两元关系“ $\geq$ ”，则从 R 中得到的严格偏好关系等价于“ $>$ ”，无差异关系等价于“ $=$ ”。

定义为两元关系的偏好概念虽然很有用，但是我们最终想研究的是行为。给定偏好关系，决策者的行为是理性的，只要他选择的结果带给他的价值至少等于其他任何结果所实现的价值。换句话说，理性决策者选择 $x^* \in X$ ，只要对每个 $y \in X$ ， $x^* Ry$ 。但是，如果不对偏好施加更多的限制条件，我们就无法保证存在这样的最优结果。我们来讨论 X 和 R 必须满足哪些条件，才能够保证这样的最优选择既有意义且存在。首先给出下面的正式定义。

定义 2.2 对选择集 X 上的弱偏好关系 R ，最优集 $M(R, X) \subset X$ 被定义为 $M(R, X) = \{x \in X: \text{对所有的 } y \in X, xRy\}$ [读作“ $M(R, X)$ 是由 X 中的元素 x 构

① 在数学上，集合 X 上的两元关系 R 是笛卡尔集 $X \times X$ 的子集。如果 $(x, y) \in R$ ，则 xRy 。

成的集合,只要对 X 中的所有元素 y , xRy ”]。

理性概念的一个基本结论是,人们从最优集中选择结果。当然了,只有在最优集中有至少一个结果时,这一结论才有意义。我们来探讨在偏好满足哪些性质时, $M(R, X)$ 是非空集。

使最优集为空集的最简单的一种情况是,在一对结果之间, R 没有做出比较。如果没有 xRy 也没有 yRx , 我们就不清楚在 x 和 y 之间,理性的选择是什么。要对 X 中的所有元素进行偏好排序,需要增加两个条件。

定义 2.3 X 上的两元关系 R :

- (1) 是完全的,如果对所有的 $x, y \in X$ 且 $x \neq y$, xRy 或 yRx 或两者同时成立。
- (2) 是反身的,如果对所有的 $x \in X$, xRx 。

完全性使得人们能用 R 比较任意两个结果。这可能不是个很有争议的假设。但是,我们知道,在面临选择时,人们有时无所适从。^①反身性则是个技术性更强的假设。一些人使用的完全性定义与这里的定义略有差异,他们的完全性中包含了反身性。^②

完全性和反身性排除了无法对任意两个结果进行比较的情况,但是,它们不能保证理性选择的存在。我们还必须排除这样一种情况: xPy 且 yPz 且 zPx 。在这种情况下,不存在理性选择——在你选择 x 时,为什么要选择 y ;在你选择 z 时,为什么要选择 x ;在你选择 y 时,为什么要选择 z ? 对偏好施加以下任何一种限制,就能解决这一问题。

定义 2.4 X 上的两元关系 R :

- (1) 是传递的,如果对所有的 $x, y, z \in X$, xRy 和 yRz 意味着 xRz ;
- (2) 是准传递的,如果对所有的 $x, y, z \in X$, xPy 和 yPz 意味着 xPz ;
- (3) 是不循环的,如果在任何一个有限子集 $\{x_1, x_2, \dots, x_n\} \subseteq X$ 上,对所有的 $i < n$, x_iPx_{i+1} 意味着 x_1Rx_n 。

这些定义之间存在着微妙差别。传递性和准传递性似乎相对容易理解,但是,它们实际上是非常强的假设,甚至可能被相当合理的行为违背。例如,假设 X 是由 1 000 瓶啤酒构成的集合。在第一瓶啤酒 b_1 中,我们用一滴自来水替换一滴啤酒;在第二瓶啤酒 b_2 中,我们用两滴自来水替换两滴啤酒;如此等等,直至第 1 000 瓶啤酒 b_{1000} 。除非你是品酒大师,否则就有 $b_1Ib_2, b_2Ib_3, \dots, b_{999}Ib_{1000}$ 。由于(根据 I 的定义) xIy 意味着 xRy , 所以, $b_{1000}Rb_{999}, \dots, b_2Rb_1$ 。如果这一关系是传递的,则有 $b_{1000}Rb_1$ 。但是,显而易见的事实是 b_1Pb_{1000} 。^③然而,无循环性假设就没有这一问题,并且一般情况下足以满足我们的需要。另一方面,尽管传递性假设存在着上述问题,我们依然用它(而不是无循环性假设),这样做能简化下面许多结论。完全性、反身性和传递性一起构成了弱序概念的基础。

① 很多经济学家和心理学家在探讨完全性假设。不使用这一条件的选择理论已经出现。

② 即对所有的 $x, y \in X$, xRy 或 yRx 或两者同时成立。

③ 这类似于 Guinness 牌子的啤酒和 Coors Light 牌子的啤酒之间的差别。

定义 2.5 给定集合 X , 弱序是具有完全性、反身性和传递性的两元关系。

实数集上的两元关系“ $\geq$ ”满足弱序的所有条件。现在给出本章的第一个结论。

定理 2.1 如果 X 是有限集且 R 是弱序, 则 $M(R, X) \neq \emptyset$ 。

这个定理指出, 只要选择集是有限集且 R 是完全的、反身的和传递的, 就存在最优选择。

证明: 设 X 是有限集且 R 是完全的、反身的和传递的。我们在 X 中的元素数量上用归纳法证明这一结论(见数学附录中的数学归纳法部分)。

第 1 步: 设 X 中只有一个元素, 即 $X = \{x\}$ 。根据反身性, xRx , 于是, $M(R, X) = \{x\}$ 。

第 2 步: 我们证明, 如果对包含 n 个元素的任何集合 X' 和 X' 上的弱序 R' , 此命题为真, 则对包含 $n+1$ 个元素的任何集合 X 和 X 上的弱序 R , 此命题为真。

第 2 步的证明: 设对包含 n 个元素的任意集合 X' 和 X' 上的弱序 R' , $M(R', X') \neq \emptyset$ 。设集合 X 包含 $n+1$ 个元素, R 为 X 上的弱序。对任意的 $x \in X$, $X = X' \cup \{x\}$, 其中 X' 是包含 n 个元素的集合。设 R' 是 R 在 X' 上的限制(即 $R' = R \cap (X' \times X')$)。根据假设, $M(R', X') \neq \emptyset$ 。于是, 对任意的 $y \in M(R', X')$, 根据完全性假设, yRx 或 xRy 或二者同时成立。

如果 yRx , 则对所有的 $z \in X' \cup \{x\}$, yRz 。于是, $y \in M(R, X)$ 。第 2 步中的结论因此得到证明。如果 xRy , 由于 $y \in M(R', X')$, 即对任意的 $z \in X'$, yRz , 于是, 对任何 $z \in X'$, xRy 且 yRz 。 R 具有传递性, 所以, 对任意的 $z \in X'$, xRz 。就是说, 对任意的 $w \in X$, xRw 。因此, $x \in M(R, X)$ 。第 2 步中的结论得到了证明。

根据数学归纳法, 第 1 步和第 2 步相结合, 证明了这一定理。□

弱序实际上不是 $M(R, X)$ 非空的必要条件。我们给出一个更一般的结论, 只是它的证明有些复杂, 我们将其留做习题。

定理 2.2 X 是有限集, R 是 X 上的完全的和反身的两元关系。对任意非空子集 $S \subset X$, $M(R, S) \neq \emptyset$, 当且仅当 R 是不循环的。

即使在选择空间为有限集且无不确定性的简单情况中, 选择理论依然相当丰富。想在这方面有更深入的了解, Austen-Smith 和 Banks(1999)是本相当好的参考书。在技术性更强的下一节中, 我们讨论结果空间为非有限集时的理性选择。对这种选择集, 我们将得到类似定理 2.1 的结论。它们在概念上类似, 但是, 后者要求对选择集和偏好施加更多的数学结构。

2.2 连续选择空间*

2.2.1 $M(R, X)$ 的非空性

在定理 2.1 的证明中, 有限选择空间的假定是关键, 它使我们能用数学归纳法

进行证明。但是,在有无限个选项时,这一方法行不通。如果决策者从连续统中(例如,从实数集 $\mathbb{R}$ 中或者从集合 $[0, 1] = \{x \in \mathbb{R} : x \geq 0 \text{ 且 } x \leq 1\}$ 中)做出选择,就需要对偏好施加更多限制,借此保证 $M(R, X) \neq \emptyset$ 。我们用两个简单例子说明这一点。

例题 2.1 设 $X = (0, 1)$ (或 $X = \mathbb{R}$)并设 R 与实数集上的两元关系“ $\geq$ ”等价,即 xRy 当且仅当 $x \geq y$ 。则集合 $M(\geq, X)$ 是空的。

我们来分析 $M(\geq, (0, 1))$ 为什么是空集。对每个 $x \in X$,存在 $y \in X$ 使得 $y > x$ 。就是说,对所有的 $y \in X$,不存在 x 使得 xRy 。集合 $(0, 1)$ 中没有最大元素——这一事实是这一例子的关键。但是, X 如果是闭区间,例如 $X = [0, 1]$,就不会出现这种问题: $M(\geq, [0, 1]) = \{1\}$ 。这个例子告诉我们,最优集的非空性可能取决于选择集是否“闭的”。下面的例子则阐述了另一种情况。

例题 2.2 设 $X = [0, 1]$ 并定义 R 为:如果 $\max\{x, y\} \leq \frac{1}{2}$ 且 $x \geq y$ 或 $\min\{x, y\} > \frac{1}{2}$ 且 $x \leq y$ 或 $x > \frac{1}{2}$ 且 $y \leq \frac{1}{2}$,则 xRy 。则集合 $M(R, X)$ 是空的。

在 $[0, \frac{1}{2}]$ 中没有元素属于 $M(R, X)$ —— $[0, \frac{1}{2}]$ 中的任何元素都劣于 $(\frac{1}{2}, 1]$ 中的每个元素。 $(\frac{1}{2}, 1]$ 中的元素也不属于 $M(R, X)$,因为在这一区间中,越接近 $\frac{1}{2}$ 的点,越被决策者偏好,但 $\frac{1}{2}$ 却不在这个集合中。因此,这个例子的结果与第一个例子相同。但在这里,问题不是出在选择集 X 上—— X 是闭区间,而是出在 R 上。 R 在 $\frac{1}{2}$ 处发生跳跃。略小于或等于 $\frac{1}{2}$ 的结果,是决策者最不偏好的结果;略大于 $\frac{1}{2}$ 的结果,是决策者最偏好的结果。偏好上的这种不连续性导致最优集是空集。

在把这些例子和我们从这些例子中得到的直观认识转化为一般性结论之前,我们介绍几个数学概念。^①假定偏好被定义在 n 维欧几里得空间上,集合 $X \subset \mathbb{R}^n$ 是选择集。这种空间中的点可以表示为向量 $x = (x^1, x^2, \dots, x^n)$,其中的坐标 x^i 是 $\mathbb{R}$ 中的点。

我们首先要知道的是,集合 X 是开集还是闭集。我们用 $\mathbb{R}$ 中的例子说明什么是开集。集合 $A \subseteq \mathbb{R}$ 是开的,如果对每个点 $x \in A$,存在某个实数 $\varepsilon > 0$ 使得对所有的 $y \in X$,如果 $|x - y| < \varepsilon$,则 $y \in A$ 。就是说,一个集合是开的,如果这个集合中任意一点附近的所有的点都是这个集合中的元素。集合 $(0, 1)$ 显然是开集。给定 $(0, 1)$ 中的任意一点,在这个集合中存在着比它更大的点同时也存在着比它更

① 准确地说,我们将介绍的是一些拓扑概念。想进一步学习选择理论的读者,应仔细学习本书的数学附录;更好的办法是找本实分析教材进行学习。Gaughan(1993)是本难度适中的入门教材, A. N. Kolmogorov 和 S. V. Fomin(1970)则更全面。

小的点。换言之,对任意的点 $x \in (0, 1)$, 存在实数 $\varepsilon > 0$ 使得 $x - \varepsilon$ 和 $x + \varepsilon$ 也在 $(0, 1)$ 内。一个集合是闭的, 如果它的补集是开集。 $(0, 1)$ 是开集, 所以, $(-\infty, 0] \cup [1, \infty)$ 是闭集。类似 $[0, 1]$ 的区间是闭集。还有一些集合既不是开集也不是闭集, 如集合 $[0, 1)$ 。

要把这些概念推广到 n 维欧几里得空间中, 我们需要用到范数的概念, 用范数衡量两个点之间的距离:

$$\|x - y\| = \sqrt{\sum_{i=1}^n (x^i - y^i)^2}$$

在这里, $\|x - y\|$ 是点 x 和点 y 之间的距离, 是对 $\mathbb{R}$ 中绝对值概念的推广。有了对距离的这一定义, 我们就能把区间推广为球。

定义 2.6 以 $x \in X$ 为圆心, 以 $\varepsilon > 0$ 为半径的开球为集合 $B(x, \varepsilon) = \{y \in X: \|x - y\| < \varepsilon\}$ 。

我们来定义开集。

定义 2.7 集合 $A \subset \mathbb{R}^n$ 是开集, 如果对每个 $x \in A$, 存在某个 $\varepsilon > 0$ 使得 $B(x, \varepsilon) \subset A$ 。

如前所述, 一个集合是闭集, 当且仅当它的补集是开集。闭集因此具有的一个特征是, 其中一些点位于其边界上, 边界点的所有的开球都包含不属于这个集合的点。

定义 2.8 集合 $A \subset \mathbb{R}^n$ 是闭的, 如果它的补 $C = \mathbb{R}^n \setminus A$ 是开集。

在第一个例子中, X 是开集, 所以围绕着 X 中的每个点 x , 都有开球包含于 X 中。在每个开球中, 都有弱优于 x 的点, 所以不存在最优集。如果 $X = [0, 1]$, 则围绕着 $x = 1$ 的每个开球都包含不属于 X 的点。由于优于 $x = 1$ 的所有的点都不在 $[0, 1]$ 中, 所以, $M(\geq, [0, 1]) = 1$ 。但是, 选择集为闭集并不是最优集为非空集的充分条件。 $(-\infty, 0] \cup [1, \infty)$ 是闭集, 但是, $M(\geq, (-\infty, 0] \cup [1, \infty))$ 都是空集。原因是, 这一集合没有上界, 所以对任意的 x , 存在 $y > x$ 使得 yPx 。所以, 我们还需要另一重要条件: 选择集的有界性。

定义 2.9 集合 $A \subset \mathbb{R}^n$ 是有界集, 如果存在某个实数 b 使得对每个 $x \in A$, $\|x\| < b$ 。

集合 $(-\infty, 0] \cup [1, \infty)$ 显然不满足这一标准, 所以, 通过要求选择集是有界集, 就能把这一集合排除掉。从例题 2.1 中容易看出, 只要 X 是闭的和有界的, $M(\geq, X)$ 就是非空的。在 $\mathbb{R}^n$ 中, 我们常用到下面的定义。

定义 2.10 集合 $A \subset \mathbb{R}^n$ 是紧的, 如果它是闭的和有界的。

在本书全部的例题和习题中, 我们用到的集合都是欧几里得空间的子集, 所以, 对数学概念的介绍完全可以到此为止。在任意选择空间中, 紧性与闭性和有界性之间的等价关系并不成立。有意思的是, 利用紧集的更一般的定义, 最优集的非空性的证明却更加容易(但是我们从例题中获得的一些直观认识却失去了用武之地)。紧集的更一般的定义, 建立在被称为开覆盖的集合族基础上。集合 A 的一个

开覆盖,是一个开集族,它的并包含 A 。

定义 2.11 给定集合 A , A 的一个开覆盖是一个集合族 $\{O_\theta\}_{\theta \in \Theta}$, 其中的 Θ 是任意的指标集, 对每个 $\theta \in \Theta$, 集合 O_θ 是开的, 且 $A \subset \bigcup_{\theta \in \Theta} O_\theta$ (换言之, 如果 $x \in A$, 则存在某个 $\theta \in \Theta$ 使得 $x \in O_\theta$)。

我们给出紧集的一般定义:

定义 2.12 集合 A 是紧的, 如果对 A 的每个开覆盖 $\{O_\theta\}_{\theta \in \Theta}$, 都存在有限集 $B \subset \Theta$, 使得有限覆盖 $\{O_\theta\}_{\theta \in B}$ 是 A 的覆盖 (即 $A \subset \bigcup_{\theta \in B} O_\theta$)。

对不熟悉实分析的读者, 上述定义不免有些费解。我们用 $\mathbb{R}$ 和 $\mathbb{R}$ 的子集 $[0, 1]$ 和 $(0, 1)$ 来阐述这些定义。我们已经知道 $(0, 1)$ 不是紧集, 因为它不是闭集。我们现在用定义 2.12 说明 $(0, 1)$ 为什么不是紧集。首先构造集合 $(0, 1)$ 的一个开覆盖: 对每个 $\theta \in \Theta = \{3, 4, 5, \dots\}$, 定义 $O_\theta = \left(\frac{1}{\theta}, 1 - \frac{1}{\theta}\right)$ 。这是一个以 $\frac{1}{2}$ 为中心的开区间族。随着 θ 趋向无穷大, 区间长度趋向 1。那么, $\{O_\theta\}_{\theta \in \Theta}$ 是 $(0, 1)$ 的开覆盖吗? 答案是肯定的。对每个 $x \in (0, 1)$, 存在足够大的 θ 使得 $x \in \left(\frac{1}{\theta}, 1 - \frac{1}{\theta}\right)$ 。换句话说, 我们确实构造了 $(0, 1)$ 的一个开覆盖。如果 $(0, 1)$ 是紧集, 根据定义 2.12 应存在有限子集 $B \subset \{3, 4, 5, \dots\}$ 使得 $(0, 1) \subset \bigcup_{\theta \in B} O_\theta$ 。但是, 在任何一个有限子集 B 中, 都有一个最大元素 $\theta^* \in B$ ①, $\frac{1}{\theta^*}$ 严格大于 0。由于 $(0, 1)$ 包含着任意地接近 0 的点, 所以, $\frac{1}{\theta^*}$ 严格大于 $(0, 1)$ 中某个元素。就是说, 对上述开覆盖的任何一个有限子集, 对任意的 $\theta \in B$, 我们都能找到 $(0, 1)$ 中的元素, 它不包含于集合 O_θ 。因此, 根据定义 2.12, $(0, 1)$ 不是紧集。你不妨用开覆盖的定义, 证明 $[0, 1]$ 是紧集。②

在明确了 X 应具有的特征后, 我们再来讨论 R 应具有的特征。例题 2.2 指出, R 不能有跳跃。给定任意一点 $x \in \mathbb{R}^n$, 我们用点 x 的上半集 (upper contour set) 和下半集 (lower contour set) 来定义 R 的连续性。③

定义 2.13 给定 $\mathbb{R}^n$ 上的两元关系 R , 点 $x \in \mathbb{R}^n$ 的严格上半集为 $P(x) \equiv \{y \in \mathbb{R}^n: yPx\}$, 严格下半集为 $P^{-1}(x) \equiv \{y \in \mathbb{R}^n: xPy\}$ 。点 x 的无差异集是对决策者而言与 x 点无差异的所有的点构成的集合, 即 $I(x) \equiv \{y \in \mathbb{R}^n: yRx \text{ 且 } xRy\}$ 。

- ① 你可以证明这句话成立。首先要看到, $\geq$ 是弱序; 然后再应用前面得到的, 在有限集情况中, 最大集具有非空性的有关结论。
- ② 这里得使用 Heine-Borel 定理, 它把紧性的拓扑开覆盖的定义, 与 $\mathbb{R}^n$ 中紧集是闭集和有界集的定义, 联系在一起。对 $\mathbb{R}$ 中的这一结论, Gaughan(1993) 给出了相当详细的证明。
- ③ 在政治科学中, 点 x 的上半集 (upper contour set) 常常被称为“优于集” (preferred to set)。但是, Keith Krehbiel 多次向本书的两位作者指出, 这一术语 (以及其他许多术语) 在语义上重复。他和我们都建议读者使用我们所偏好的术语: 所偏好的集合 (preferred set)。

对任意的 x , 其严格上半集由严格优于点 x 的一切点构成; 其严格下半集由严格劣于点 x 的一切点构成; 点 x 的无差异集由与点 x 无差异的一切点构成。

定义 2.14 $\mathbb{R}^n$ 上的两元关系 R 是:

- (1) 上连续的, 如果对所有的 $x \in \mathbb{R}^n$, $P(x)$ 是开集;
- (2) 下连续的, 如果对所有的 $x \in \mathbb{R}^n$, $P^{-1}(x)$ 是开集;
- (3) 连续的, 如果它既是下连续的又是上连续的。

我们来分析这些定义的含义。给定点 x , 根据完全性, 任何一个点 y 或者是 $P(x)$ 中的元素, 或者是 $P^{-1}(x)$ 的元素, 或者是 $I(x)$ 的元素。根据连续性, 如果 $y \in P(x)$, 则足够地接近 y 的所有的点亦在 $P(x)$ 中。如果 $y \in P^{-1}(x)$, 则 y 附近的点亦在 $P^{-1}(x)$ 中。就是说, y 的微小扰动不影响它与 x 之间的偏好关系。

在例题 2.2 中, 我们能看到连续性是如何排除掉异常行为的。在这个例子中, 劣于 $\frac{1}{2} + \varepsilon$ 的集合为 $P^{-1}\left(\frac{1}{2} + \varepsilon\right) = \left(-\infty, \frac{1}{2}\right] \cup \left(\frac{1}{2} + \varepsilon, 1\right]$, 它不是开集。如果偏好是下连续的, 则在偏好序上不应该有任何跳跃。

现在给出最优集为非空集的充分条件。

定理 2.3 如果 X 是非空的和紧的, X 上的两元关系 R 是完全的、反身的、传递的和下连续的, 则 $M(R, X) \neq \emptyset$ 。

和本书大部分章节相比, 这一结论的证明中包含的技术性更强。但是, 这一结论在相当一般的背景中成立。这使得我们能够用它分析各种选择问题, 无论在这些个问题中, x 是结果构成的无穷数列, 是函数还是概率分布。

证明: 设 X 是非空的和紧的, R 是完全的、反身的、传递的和下连续的。我们使用反证法, 设 $M(R, X) = \emptyset$ 。于是, 对 X 中的任何一个点, 存在某个 $\alpha \in X$ 使得此点包含在 $P^{-1}(\alpha)$ 中。 R 是下连续的, 因此, 每个 $P^{-1}(\alpha)$ 都是开集, $\{P^{-1}(\alpha)\}_{\alpha \in X}$ 是 X 的开覆盖。 X 是紧集, 所以, 存在有限子集 $B \subset X$ 使得集合族 $\{P^{-1}(\alpha)\}_{\alpha \in B}$ 是 X 的开覆盖。对任何一个 $x \in X$, 存在某个 $\alpha \in B$ (α 随 x 而不同) 使得 $x \in P^{-1}(\alpha)$ 。 B 是有限集, R 是完全的、反身的和传递的, 于是根据定理 2.1, $M(R, B) \neq \emptyset$, 就是说, 存在 $x^* \in M(R, B)$ 。给定任意的 $y \in X$, 它或者是 $M(R, B)$ 中的元素或者不是。如果 $y \in M(R, B)$, 根据定义, $x^* I y$, 于是 $x^* R y$ 。如果 $y \notin M(R, B)$, 由于 $\{P^{-1}(\alpha)\}_{\alpha \in B}$ 覆盖 X , 因此存在某个 $\alpha \in B$ 使得 $y \in P^{-1}(\alpha)$, 即 $\alpha R y$ 。而 $x^* \in M(R, B)$, 因此 $x^* R \alpha$ 。又由于 X 上的两元关系 R 具有传递性, 所以, $x^* R y$ 。于是, 对所有的 $y \in X$, $x^* R y$ 。也就是说, $x^* \in M(R, X)$, 而这与 $M(R, X)$ 是空集的假定矛盾。□

定理 2.3 给出的是最优集为非空集的充分条件而不是必要条件。在我们所分析的问题中, 有时候 X 可能是无界的或者不是闭集, R 是不连续的。此时, 我们必须用其他方法判断 $M(R, X)$ 的非空性。 X 如果不是紧集, 一般就需要对 R 施加更强的假设; R 如果不是连续的, 则需要对 X 施加更强的假设。

2.2.2 $M(R, X)$ 的唯一性

我们需要知道 $M(R, X)$ 中是否只有一个元素。如果选择集为有限集, 则通过假定所有偏好都是严格偏好, 就能保证 $M(R, X)$ 中只有一个元素。这一假定实际上排除了结果之间无差异的情况, 于是在 X 为有限集时, $M(R, X)$ 中的元素不可能超过一个。

但是, 如果选择空间不是有限集, 我们就需要增加条件来保证 $M(R, X)$ 中只有一个元素。许多应用研究同时对 X 和 R 施加限制条件。我们一般假定 X 是凸集。在 X 为凸集时, 如果 x 和 y 是 X 中的点, 则 x 和 y 之间的线段上的所有的点肯定在 X 中。

定义 2.15 集合 $X \subset \mathbb{R}^n$ 是凸的, 如果对任意的 $x, y \in X$, 对每个 $\lambda \in [0, 1]$, $\lambda x + (1-\lambda)y$ 是 X 中的元素。

点 $\lambda x + (1-\lambda)y$ 被称为 x 和 y 的凸组合(或加权平均)。集合 $[0, 1]$ 是凸集, 因为 $[0, 1]$ 中任意两个点之间的点, 也在 $[0, 1]$ 中。集合 $X = [0, \frac{1}{4}] \cup [\frac{3}{4}, 1]$ 不是凸集, 因为对任意的 $\lambda \in (0, 1)$, $\lambda \frac{1}{4} + (1-\lambda) \frac{3}{4} \notin X$ 。直观地看, 凸性条件要求结果集中没有“窟窿”。如果结果集的维数超过一维, 凸性条件还要求结果集的表面没有“皱褶”。看看你的手掌。大拇指上的点和食指上的点构成的凸组合, 不在手掌上。^①所以, 你的手掌不是凸的。

我们同时施加于偏好上的充分条件也被称为凸性。

定义 2.16 定义在凸集 X 上的偏好关系 R 是凸的, 如果 xRy 意味着对任意的 $\lambda \in (0, 1)$ 和所有不同的点 $x, y \in X$, $[\lambda x + (1-\lambda)y]Ry$ 。偏好关系 R 是严格凸的, 如果对任意的 $\lambda \in (0, 1)$ 和所有不同的点 $x, y \in X$, $[\lambda x + (1-\lambda)y]Py$ 。

凸偏好的性质是, 如果在 x 和 y 之间, 决策者偏好 x , 则在 x 和 y 的凸组合与 y 之间, 他偏好凸组合。严格凸性则更进一步: 即使决策者在 x 和 y 之间无差异, 相对于 x 或 y , 他仍偏好 x 和 y 的凸组合。在本章的习题中, 我们请你证明, R 的凸性意味着严格优于点 x 的集合 $P(x)$ 是凸集。对任何一点 x , 由于严格优于 x 的集合是凸集, 所以这种集合没有“窟窿”或“皱褶”。严格凸性则保证在这类集合的边界上没有平直部分。

下面的结论相对容易证明。

定理 2.4 如果 X 是凸的且(定义在 X 上的) R 是严格凸的, 则 $M(R, X)$ 包含至多一个元素。

证明: 使用反证法。设 X 是凸集, R 是严格凸的, x 和 y 是在 $M(R, X)$ 中两个

① 博弈论学家花费了大量时间探讨选择集的这种非凸性。

不同的点。对任意的 $\lambda \in (0, 1)$, 由于 X 是凸的, 所以 $\lambda x + (1 - \lambda)y$ 在 X 中。由于 R 是严格凸的, 所以 $[\lambda x + (1 - \lambda)y]Py$ 。这与 $y \in M(R, X)$ 的假设矛盾。□

在选择集是紧集且弱序具有下连续性时, 定理 2.3 保证了理性选择的存在。如果选择集还是凸集且偏好序是严格凸的, 则理性选择具有唯一性。

2.3 效用理论

前述选择和理性模型使用的是两元偏好和最优集的概念。除非是最简单的模型, 否则, 使用两元关系不是件易事。与此相反, 数字容易处理。如果我们能赋予结果集中每个元素一个数字, 我们就能用运算符 $\geq$ 比较各个结果。本节探讨在什么条件下, 我们能够用实数表示结果, 并用 $\geq$ 表示偏好。换言之, 我们想用效用函数(定义域为 X 的实值函数)表示偏好从而使得:

$$u(x) \geq u(y) \text{ 意味着 } xRy$$

$$u(x) > u(y) \text{ 意味着 } xPy$$

$$u(x) = u(y) \text{ 意味着 } xIy$$

效用概念是在过去 300 年中哲学和道德讨论的主题, 但在这里, 我们使用的是效用的一个狭义定义。效用只是对偏好的数字表示, 在这种表示中, $\geq$ 是偏好运算符——我们不赋予效用任何规范内涵。

在当前讨论中, 效用函数只具有序数特征: 我们用效用函数排列各个结果。它们并没有告诉我们, 一个结果优于另一个结果的程度。 $u(x) - u(y)$ 的值没有任何意义。给定任何函数 w , 如果 $w(x) \geq w(y)$ 当且仅当 $u(x) \geq u(y)$, 则这个函数和 u 表示相同的偏好。这就是说, 比较不同人的效用水平一般是无意义的。但是, 在下一章中, 关于不确定性条件下的选择的标准模型假定效用函数还包含序数特征以外的信息。

我们给出效用函数的正式定义:

定义 2.17 给定 X 和 X 上的偏好关系 R , 效用函数 $u: X \rightarrow \mathbb{R}$ 表示 R , 如果对所有的 $x, y \in X$, $u(x) \geq u(y)$ 当且仅当 xRy 。

利用这一定义, 能很容易地证明, $u(x) > u(y)$ 当且仅当 xPy , $u(x) = u(y)$ 当且仅当 xIy 。如果 X 是有限集, 则表示 R 的效用函数的存在决定于 R 是完全的、反身的和传递的关系。

我们用效用函数定义来分析决策者的最优选择。 x 是使函数 $u: X \rightarrow \mathbb{R}$ 取最大值的元素, 如果对所有的 $y \in X$, $u(x) \geq u(y)$ 。下面的结论指出, 最大值元素的存在等价于 $M(R, X)$ 为非空集。

定理 2.5 如果函数 $u(\cdot)$ 是选择集 X 上的偏好关系 R 的效用表示, 则 $M(R, X) = \arg \max_{x \in X} \{u(x)\}$ 。

证明: 首先证明 $M(R, X) \subset \arg \max_{x \in X} \{u(x)\}$ 。设 $u(\cdot)$ 表示 R , $x' \in M(R, X)$ 。由于 $x' \in M(R, X)$, 因此对所有的 $y \in X$, $x'Ry$ 。于是, 对所有的 $y \in X$,

$u(x') \geq u(y)$ 。因此, $x' \in \arg \max_{x \in X} \{u(x)\}$ 。

再来证明 $\arg \max_{x \in X} \{u(x)\} \subset M(R, X)$ 。设 $u(\cdot)$ 表示 R , $x' \in \arg \max_{x \in X} \{u(x)\}$ 。则对所有的 $y \in X$, $u(x') \geq u(y)$, 也就是说对所有的 $y \in X$, $x' R y$ 。因此, $x' \in M(R, X)$ 。□

如果 X 是有限集且 R 是完全的、反身的和传递的, 则 $M(R, X)$ 是非空集(定理 2.1); 因此, $u(x)$ 的最大值元素肯定存在。但是, 如果 X 不是有限集, 则要保证最大值元素的存在, 我们需要对 X 和效用函数施加更多限制。下一节探讨定义在非有限的结果空间上的效用函数。

2.4 连续选择空间上的效用表示*

$M(R, X)$ 如果为非空集, 则偏好的连续性是个重要条件。由于这一原因, 我们常常假定效用函数是连续的。

定义 2.18 函数 $f: X \rightarrow \mathbb{R}$ 是连续的, 如果对每个 $x \in X$, 以下陈述成立: 对每个 $\varepsilon > 0$, 存在 $\delta > 0$ 使得当 $\|x - y\| < \delta$ 时, $|f(x) - f(y)| < \varepsilon$ 。

正如在高中阶段所学到的, 连续函数是不抬起笔尖就能画出的函数。实际上, 对彼此接近的结果, 连续效用函数产生的效用几乎相同。

下面关于偏好的充分条件保证了连续的效用函数的存在。

定理 2.6 (Debreu, 1959) 如果 $X \subset \mathbb{R}^n$ 且 R 是完全的、反身的、传递的和连续的, 则存在表示 R 的连续效用函数 $u: X \rightarrow \mathbb{R}$ 。

这里不证明这一定理。^①在习题中, 我们请你证明这一定理的逆命题。我们有类似定理 2.3 的结论:

定理 2.7 如果 $X \subset \mathbb{R}^n$ 是紧的且 $u: X \rightarrow \mathbb{R}$ 是连续的, 则存在最大值元素。

这一结论有时被称为 Weierstrass 定理, 这里不证明这一定理[Royden(1988)给出了一个证明]; 定理 2.3 实际上比定理 2.7 更强, 它只要求下连续性(对每个 x , 集合 $\{y: u(y) < u(x)\}$ 是开集)和紧性。

前面说过, 效用函数带有任意性; 它们包含序数信息却无基数信息。就是说, 具体的效用值本身没有任何内在意义。对任意两个结果 $x, y \in X$, 真正重要的是 $u(x)$ 和 $u(y)$ 之间的大小关系。 $f: \mathbb{R} \rightarrow \mathbb{R}$ 是严格增函数, 如果对所有的 $x, y \in X$, $x > y$ 意味着 $f(x) > f(y)$ 。效用函数只服从严格递增变换。就是说, 如果 $u: X \rightarrow \mathbb{R}$ 表示偏好关系 R 且 $f(\cdot)$ 是严格递增变换, 则 $f \circ u: X \rightarrow \mathbb{R}$ 表示 R 。把 x 映射为 $f(u(x))$ 的函数。效用函数的任何严格递增变换都不影响最终选择, 效用函数的取值没有任何实际意义。

前面的讨论给出了效用函数有最大值的充分条件, 但是没有谈到最大值元素的特征。如果效用函数是可微的, 我们就能用微积分分析最优选择的特征。书末

① 感兴趣的读者, 可以参考 Debreu(1959)。

的数学附录介绍了微积分和优化理论中的一些重要结论。

2.5 空间偏好

在大部分经济学模型中,结果常常表示为货币(收入、财富、工资和利润等)或者商品。在这种情况下,很自然地,值越大的结果越好。换言之,经济学中探讨的许多偏好具有非满足性,决策者或者认为多多益善(例如货币)或者认为越少越好(如空气污染)。而在政治博弈论中,我们所研究的许多结果体现为政策,人们最好的结果(如税收、福利或限制堕胎的政策等)可能既不是0也不是无穷大。例如,在某一税率水平以下时,选民的效用可能随税率的上升而上升;当税率在这一水平以上时,他的效用可能随税率的上升而下降。某位选民可能希望政府限制堕胎,但是,只希望把怀孕三个月之后的堕胎行为定为非法。就是说,在政治科学研究中,我们常常需要假定政治决策者的偏好具有满足性。正式地说,某个人的偏好具有满足性,如果 $M(R, X)$ 所包含的元素位于结果空间 X 的内点集中。或者说,偏好具有满足性,如果使 $u: X \rightarrow \mathbb{R}$ 取最大值的元素在 X 的内点集中。从图 2.1 中我们能看出满足性偏好和非满足性偏好之间的区别。

满足性偏好多出现于空间模型中——这种模型把政策结果表示为 $\mathbb{R}^d$ 中的点。理论上说,我们可以假定相当一般的此类偏好,但是,在实际研究中(和在本书的大部分模型中),我们一般假定决策者有单峰的和对称的偏好。第 4 章将详细讨论单峰偏好,这里仅指出一点:单峰性意味着决策者的最优集中只有一个元素,效用函数只在一个点上取得最大值。这个最被偏好的政策结果被称为决策者的理想点。对称性假设要求决策者的效用在任何方向上以相同速度下降,就是说,偏好是政策结果和决策者的理想点之间的距离的减函数。

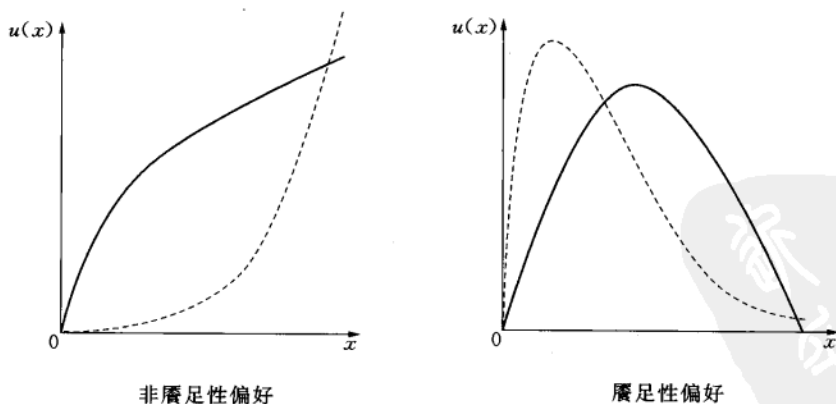


图 2.1 满足性效用函数和非满足性效用函数

如果政策空间是一维空间,则单峰对称偏好可以表示为效用函数 $u_i(x) = h(-|x - z_i|)$, 其中 z_i 是决策者 i 的理想点, h 是个增函数。在这类效用函数

中,使用得最多的是线性效用函数 $u_i(x) = -|x - z_i|$ 和二次方效用函数 $u_i(x) = -(x - z_i)^2$ 。图 2.2 给出了这两种函数的形状。

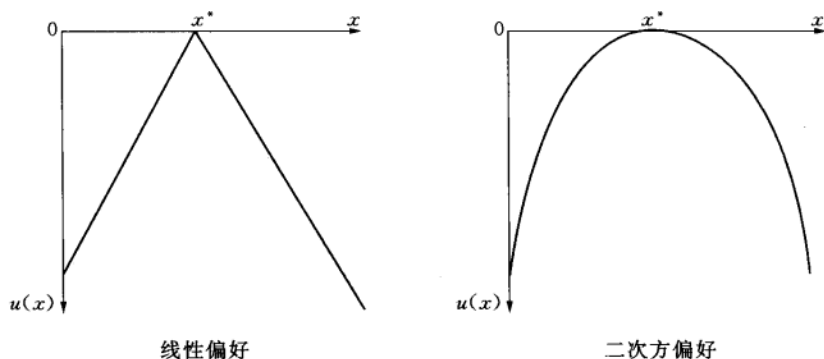


图 2.2 线性偏好和二次方偏好

在多维结果空间中(即 $X \subset \mathbb{R}^d$), 距离一般用欧几里得范数衡量:

$$\|x - z_i\| = \sqrt{\sum_{j=1}^n (x^j - z_i^j)^2}$$

对称单峰偏好的形式于是为:

$$u_i(x) = h(-\|x - z_i\|)$$

其中, z_i 是决策者 i 的理想点, h 是增函数。

多维空间上的效用函数难以用图形描绘。但是,在二维空间中,由于每个人的偏好集(即 $P(y) = \{x \in X \mid xRy\}$) 是以其理想点为圆心的圆形区域,所以,图形分析相对简单。给定政策 y , 以理想点为圆心的穿过点 y 的圆上的所有的点,都与点 y 无差异。图 2.3 中画出了这些集合。给定一条无差异曲线,在圆内部的结果和外部的结果之间,决策者偏好前者。

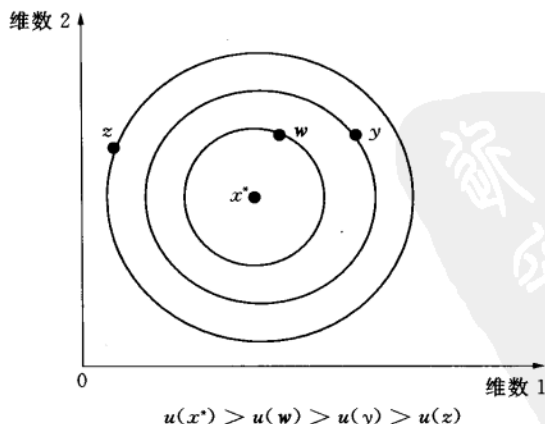


图 2.3 二维空间中二次方偏好下的无差异曲线

在政治博弈论模型中,单峰对称偏好被大量使用的一个原因是,它能帮助我们相对容易地分析决策者的选择所具有的特征。在恰当的假设下,在任何两个结果之间的选择,都可以用决策者的理想点和 $\mathbb{R}$ 中的“折点”或 $\mathbb{R}^d$ 中的“分离平面”来表述。

设决策者 i 有 $\mathbb{R}$ 上的对称单峰偏好,就是说,在 x 和 y 之间他偏好 x ,当且仅当:

$$h(-|x - z_i|) > h(-|y - z_i|)$$

如果 $x > y$,这一条件就变为:

$$z_i > c \equiv \frac{x+y}{2}$$

另一方面, yPx 当且仅当 $z_i < c$ 。因此,给定一群决策者,设结果 $x > y$,则理想点大于 x 和 y 的中点的决策者,都偏好 x ;理想点小于 x 和 y 的中点的决策者,都偏好 y 。这一性质与函数 h 没有任何关系。

这一结论适用于 $\mathbb{R}^d$ 。在 $\mathbb{R}^d$ 空间中,决策者 i 在 x 和 y 之间偏好 x ,当且仅当 $h(-\|x - z_i\|) > h(-\|y - z_i\|)$ 。我们定义分离超平面为 $C = \{c \mid \|x - c\| = \|y - c\|\}$,它等价于 $\mathbb{R}$ 中的“折点”。它把理想点分离开来:在 x 和 y 之间偏好 x 的人的理想点,和在 x 和 y 之间偏好 y 的人的理想点。只要知道决策者的理想点和 C ,我们就能探讨决策者的选择所具有的特征。

2.6 习题

习题 2.1 证明:如果 X 是有限集且 R 是 X 上的完全的和反身的两元关系,则对任意的非空集 $S \subset X$, $M(R, S) \neq \emptyset$ 当且仅当 R 是不循环的。

习题 2.2 证明:如果 X 上的两元关系 R 是传递的,则衍生自 R 的无差异关系 I 和严格偏好关系 P 是传递的。

习题 2.3 在本章中,我们首先定义了弱序 R ,然后定义 P 为: xPy 当且仅当 xRy 且无 yRx 。我们也可以首先定义严格偏好关系 P ,然后定义 R 为: xRy 当且仅当无 yPx 。

(1) 关系 P 是不对称的,如果对任何一对不同的点 $x, y \in X$, xPy 和 yPx 不同时成立。证明: P 是不对称的当且仅当 R 是完全的。

(2) 关系 P 是负传递的,如果对任意的 $x, y, z \in X$, xPy 意味着 xPz 或 zPy 或两者同时成立。证明: P 是负传递的当且仅当 R 是传递的。

习题 2.4 对下面的效用函数,找出各自对应的偏好集——你可以用图形来描述,也可以对所有的 $x \in X$,表述 $P(x) = \{y: yRx\}$ 的特征。只要可能,就请画出下列效用曲线:

(1) $u(x) = -|1 - x|$, 其中, $x \in [0, 1]$;

(2) $u(x) = -x^2$, 其中, $x \in [0, 1]$;

(3) $u(x) = \sqrt{x}$, 其中, $x \in [0, 1]$;

(4) $u(x) = -\alpha x_1^2 - (1-\alpha)x_2^2$, 其中, $x \in \mathbb{R}^2$ 。

习题 2.5 某个人在 $\mathbb{R}^d$ 上的空间偏好 R 可以表示为效用函数 $u_i(x) = h(-\|x - z_i\|)$ 。

(1) 对任意紧集 $X \subset \mathbb{R}^d$, 证明 $M(R, X)$ 非空。

(2) 分析 $M(R, X)$ 的特征。

(3) 如果 X 同时是凸集, 证明 $M(R, X)$ 中只有一个元素。

(4) 给定 $w \in \mathbb{R}^d$ 和 $k > 0$, 设 $X = \{x \in \mathbb{R}^d: \|x - w\| < k\}$ 。请把集合 $M(R, X)$ 表示为向量 z_i 与 w 和标量 k 的函数。

习题 2.6(*) 设 X 是凸集, R 是弱序。证明: 如果 R 是严格凸的, 上半集 $P(x)$ 是凸集。

习题 2.7(*) 利用定义 2.12 证明 $[0, 1]$ 是紧集。

习题 2.8(*) 设 $X \subset \mathbb{R}^n$, $u: X \rightarrow \mathbb{R}$ 是连续效用函数, 表示 X 上的两元关系 R 。证明 R 是完全的、反身的、传递的和连续的。

习题 2.9(*) 利用定理 2.7 证明 Weierstrass 定理: 如果 X 是紧的且 $u(\cdot)$ 是 X 上的连续的实值函数, 则存在点 $x \in \arg \max_{x \in X} u(x)$, 它最大化 $u(\cdot)$ 。

不确定性下的选择

前一章假定决策者能完全预测到其行动所导致的后果。本章则放松这一假设。决策者只知道各种结果以某个概率产生自他们对行动的选择——知道某些行动会增加或者降低某些结果的发生概率。换句话说,人们知道哪些行动更可能导致某些结果。在前一章的例子中, $A = \{\text{出兵, 谈判, 不作为}\}$, $X = \{\text{大胜, 小胜, 现状}\}$ 。决策者可能知道,相对于谈判,如果出兵的话,对方做出巨大让步的可能性更大。于是,在决策时,他就会把得到更好结果的可能性与各种行动的成本相权衡。只有在出兵决策更可能迫使对方做出更大让步时,且只有在对方的进一步退让十分重要时,而且在出兵成本很低时,出兵才是理性决策。在关于不确定性条件下的选择的经典理论中,这些都是基本的权衡因素。

在这种不确定性模型中,有两个基本的构成要素。首先是概率推断(beliefs),即与具体行动相联系的结果空间上的概率分布或赌局。其次是各个结果所实现的收益。这两个要素相结合,便有了定义在行动集上的 von Neumann-Morgenstern 效用函数——以古典决策理论的两位开创者的名字命名。我们将看到,与第 2 章所讨论的序数函数相比, von Neumann-Morgenstern 所依托的效用概念更强。

3.1 有限行动集和有限结果集

首先探讨行动集和结果集都为有限集的情况。和前一章一样,我们把可行的行动和结果表示为集合 $A = \{a_1, \dots, a_I\}$ 和 $X = \{x_1, \dots, x_J\}$ 。在这里,行动和结果是以某种概率分布而不是以确定的方式联系在一起。为说明这种联系,假定结果既决定于所采取的行动亦决定于自然状态 s 。从决策者角度看, s 是随机变量,例如选举日下雨或战争中导弹的命中率。在决策理论模型中(或者在决策者随机地选择行动的博弈论模型中), s 可以为其他决策者的行动。状态集为 $S = \{s_1, \dots, s_K\}$ 。决策者对每种状态的发生概率有自己的推断,我们将其表示为概率函数 $\pi(s_k) \equiv \pi_k$ 。这些概率推断必须满足概率论的基本公理,即必须在 0 和 1 之间,且其和必须等于 1。正式地说,

对每个 k , $0 \leq \pi_k \leq 1$; $\pi_1 + \pi_2 + \cdots + \pi_K = \sum_{k=1}^K \pi_k = 1$

结果函数 $\chi(a, s): A \times S \rightarrow X$ 则把行动和状态与结果联系起来。^①在表 3.1 中, 对应着状态和行动的每个组合, 有一个结果。

表 3.1

$A \backslash S$	s_1	s_2	s_3
a_1	x_1	x_2	x_2
a_2	x_1	x_2	x_3

在表 3.1 中, 无论选择哪个行动, 只要状态为 s_1 , 就有结果 x_1 。在状态 s_2 和 s_3 中, 结果决定于决策者所选择的行动。对决策者来说, 真正重要的不是状态, 而是每个行动所导致的结果的发生概率。由于在选择 a_i 时, 决策者并不知道状态是什么, 所以得到结果 x 的概率便是状态取值为 s 从而使得

$\chi(a_i, s) = x$ 的概率。设 p_{ij} 是在采取行动 a_i 后结果 x_j 的发生概率。

从表 3.1 中, 我们很容易算出这些概率: $p_{11} = \pi_1 + \pi_2$, $p_{12} = \pi_3$, $p_{13} = 0$, $p_{21} = \pi_1$, $p_{22} = \pi_2$, $p_{23} = \pi_3$ 。这些概率的一般计算公式是:

$$p_{ij} = \sum_{\{k: \chi(a_i, s_k) = x_j\}} \pi_k$$

p_{ij} 从 π_k 继承了以下性质:

对每个 i 和 j , $0 \leq p_{ij} \leq 1$;

$$\text{对每个 } i, p_{i1} + p_{i2} + \cdots + p_{iJ} = \sum_{j=1}^J p_{ij} = 1$$

我们于是可以忽略结果的发生概率对 s 的依赖, 而只使用 p_{ij} , 借此简化符号。但是在后面一些章节中, 我们依然用行动—状态来表述决策者的问题。

为简化符号, 我们定义向量 $p_i = (p_{i1}, \cdots, p_{iJ})$, 它是采取行动 a_i 时结果上的一个概率分布。由于行动和由行动产生的概率分布之间的这种对应关系, 我们可以说决策者选择行动 a_i , 或者等价地说, 决策者选择概率分布 p_i 。设 P 是所有概率分布构成的集合。由于有 J 种可能的结果, 所以, 集合 P 由有 J 个坐标的向量构成——这些坐标必须满足前述条件(每个坐标在 0 和 1 之间取值, 所有坐标值之和为 1)。我们把这一集合表示为 Δ^J , 称为 J 维单纯型。^②二维情况单纯型是 $(0, 1)$ 到 $(1, 0)$ 之间的直线, 见图 3.1 中的第一幅图。三维单纯型是顶点分别为 $(1, 0, 0)$ 、 $(0, 1, 0)$ 和 $(0, 0, 1)$ 的三角区域, 见图 3.1 中的第二幅图。

我们也可以使用图 3.2 中的树状图表示概率分布。从初始节点开始, 每个树枝对应一个具体结果, 树枝上标出了这一结果的发生概率。和概率分布 q 相比, 在概率分布 p 中, x_1 的发生概率更大, x_3 的发生概率更小, x_2 的发生概率相同。为给

① 由于我们所探讨的结果都以 a 和 s 的某种组合为前提, 所以, 我们要求行动—状态组合的数量不能小于结果数, 换言之, 我们要求 $I \cdot K \geq J$ 。

② 不熟悉向量和坐标系的读者, 不妨参考书末的数学附录。

后面的讨论打好铺垫,我们来分析决策者如何在 p 和 q 之间做出选择。首先,决策者不会根据他对 x_2 的偏好而做出选择——在这两个概率分布中,结果 x_2 的发生概率相同。这两个概率分布只在 x_1 和 x_3 的发生概率上彼此不同,所以,只有在 $x_1 R x_3$ 时,理性决策者才会选择 p 。基于这种直观认识(我们马上就给出其正式表述),我们的结论是,如果 $x_1 P x_3$, 决策者选择 p ; 如果 $x_3 P x_1$, 决策者选择 q ; 如果 $x_1 I x_3$, 决策者在这两个概率分布之间无差异。

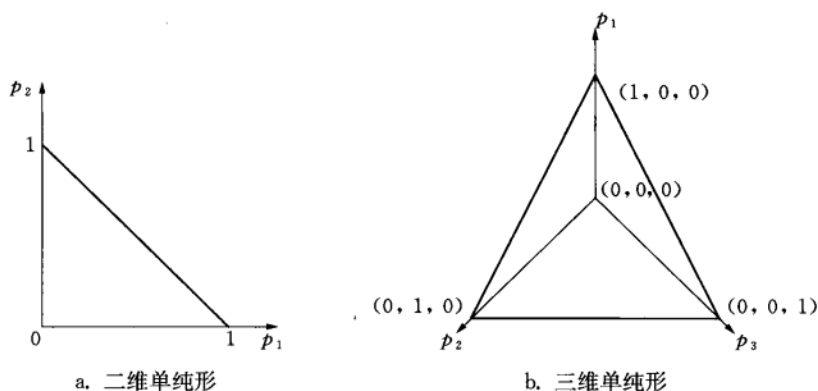


图 3.1 单纯形

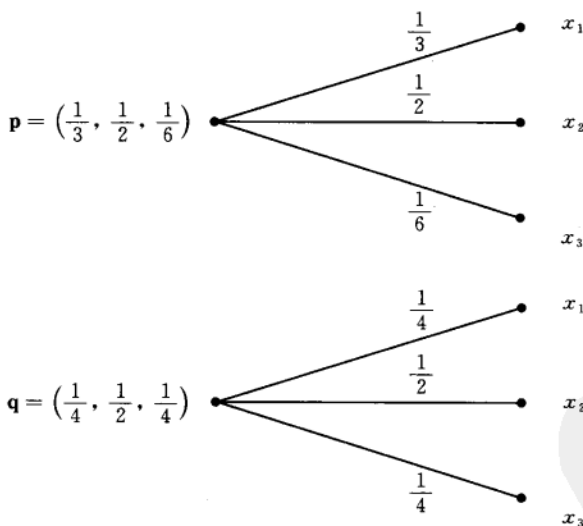


图 3.2 概率分布的树状图

在这个例子中,概率分布是简单分布,正是这一特点帮助我们得到了上面的直观结论。在简单分布中,每个结果对应着一个概率。但我们常常需要探讨更复杂的情况,其中,决策者得在概率分布上的概率分布上的概率分布……之间做出选择。这种概率分布被称为复合概率分布(compound lotteries)。P 上的复合概率分

布可表示为 $\{a_1, \dots, a_I\}$, a_i 为概率分布 p_i 的发生概率。假设一位理性决策者对着一个复合概率分布, 在这一分布中, 他以概率 $\frac{1}{4}$ 得到概率分布 p , 以概率 $\frac{3}{4}$ 得到概率分布 q 。我们可以把这一概率分布写为 $r = \frac{1}{4}p + \frac{3}{4}q$ 。图 3.3 给出了它的树状图。那么, 决策者如何在 p 、 q 和 r 之间做出选择呢? 首先, r 的出现不应该改变决策者对 p 和 q 的偏好, 就是说, 这里只需把 r 与 p 进行比较和把 r 与 q 进行比较。由于 r 的复合结构, 这种比较似乎有些难度。实际情况并非如此。理性决策者只关心与每个结果相对应的概率, 而不在乎到达各个结果的路径。因此, 他会算出得到结果 x_1 的概率为得到概率分布 p 的概率 $\frac{1}{4}$ 乘以 $\frac{1}{3}$, 加上得到概率分布 q 的概率 $\frac{3}{4}$ 乘以 $\frac{1}{4}$ 。他以同样的方法算出结果 x_2 和 x_3 的发生概率。他为每个结果算

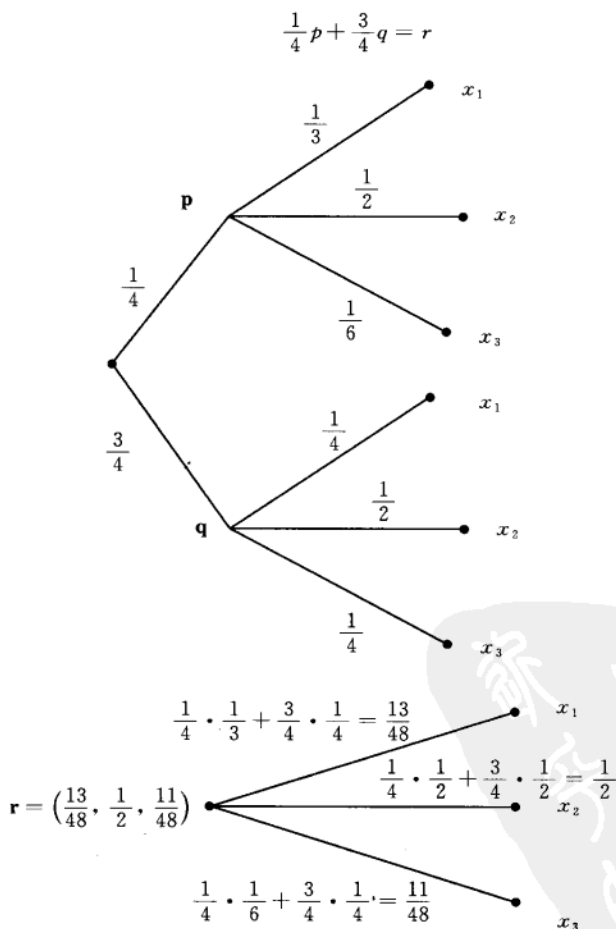


图 3.3 复合概率分布

出对应的概率,于是, r 就被表示为图 3.3 中第二幅图中的简单概率分布。正式地说, P 上的复合概率分布 $\{\alpha_1, \dots, \alpha_I\}$ ——它赋予概率分布 p_i 以概率 α_i ——都可以转化为简单概率分布,结果 x_j 的发生概率为 $\sum_{i=1}^I \alpha_i p_{ij}$ 。

在 r 被转化为简单概率分布后,决策者偏好哪个概率分布呢? 在 r 的简单形式中, x_2 的发生概率依然是 $\frac{1}{2}$,和它在 p 与 q 中的发生概率相同。所以,在比较这三个概率分布时,对 x_2 的偏好依然不重要——重要的是它们在 x_1 和 x_3 上的差异。和 q 相比,在 r 中,结果 x_1 比 x_3 更可能发生,所以偏好为 $x_1 P x_3$ 的决策者,在 r 和 q 之间偏好 r ;在 p 和 r 之间偏好 p ,因为在 p 中,结果 x_1 比 x_3 更可能发生。

我们现在知道,有结果集 X 上的偏好关系的决策者是如何比较各个概率分布的。上面的讨论揭示了不确定性条件下的选择的标准模型应具有的一些特征。这一模型定义了 P 上的弱偏好,借此分析不确定性下的选择。正如效用函数大大简化了我们对确定性条件下选择的分析,期望效用模型简化了我们对不确定性条件下选择的分析。前面例子中得到的直观认识被精炼地概括为关于 P 上的弱偏好 R 的四个公理。

公理 3.1 完全性和传递性: P 上的弱偏好关系 R 是完全的、反身的和传递的。

公理 3.2 复合概率分布的简化:对任意的 $\alpha \in [0, 1]$ 和 $p \in P$, $p I [\alpha p + (1 - \alpha)p]$ 。

公理 3.3 连续性:设 p 、 q 和 r 是 P 中的三个概率分布。取 $\alpha \in [0, 1]$, 则使得 $[\alpha p + (1 - \alpha)r] R q$ 的所有的 α 构成的集合,是闭区间。取 $\beta \in [0, 1]$, 则使得 $q R [\beta p + (1 - \beta)r]$ 的所有的 β 构成的集合,是闭区间。^①

公理 3.4 独立性:设 p 、 q 和 r 是 P 中的三个概率分布。对任意的 $\alpha \in (0, 1)$, $p R q$ 当且仅当 $[\alpha p + (1 - \alpha)r] R [\alpha q + (1 - \alpha)r]$ 。

上述公理的含义相当直白。首先,和对结果之间的比较一样,对两个概率分布,我们可以用偏好关系进行比较,并且概率分布上的偏好关系不循环。在这里,这一个公理是关键。而且,和第 2 章一样,传递性可以推广到无差异和严格偏好关系中。第二个公理是对图 3.3 的数学表述,它指出,决策者只关心各个结果的发生概率,而不在乎产生这些结果的概率过程。^②

和前两个公理相比,连续性公理更抽象。它告诉我们,结果的发生概率上的微小变化不会使偏好序发生显著变化。根据这一公理,如果 $p P q$,则足够接近 p 的所有概率分布都优于 q 。如果 $p P q$,略微地改变 p ,譬如略微提高某个很差结果的发生概率,不会影响偏好序。从连续性公理,我们能得到下面的结论。

① 数学附录对闭集有详细的讨论。就目前的目的而言,我们只需知道闭区间 $[a, b]$ 包含端点 a 和 b 和它们之间所有的点。

② 在一些教材中,在作者定义概率分布上的偏好时,这一假设只是隐含在其中。我们明确给出这一假设的目的,是强调这一理论不考虑有关概率分布的表示形式的细节特征。

引理 3.1 如果 $pRqRr$, 则存在 $\lambda \in [0, 1]$ 使得 $[\lambda p + (1-\lambda)r]Iq$ 。

证明: 设 $pRqRr$ 。对每个 $\lambda \in [0, 1]$, $[\lambda p + (1-\lambda)r]Pq$ 或 $qP[\lambda p + (1-\lambda)r]$ 或 $[\lambda p + (1-\lambda)r]Iq$ 。如果 pIq 或 rIq , 则在 $\lambda = 1$ 或 $\lambda = 0$ 时, 这一引理成立。我们还需证明在 $pPqPr$ 时, 这一引理成立。如果 $pPqPr$, 则 $\{\alpha: [\alpha p + (1-\alpha)r]Rq\}$ 和 $\{\beta: qR[\beta p + (1-\beta)r]\}$ 都是非空集。根据连续性公理, 这两个集合是闭集。因此, 第一个集合有最小元素, 设为 $\underline{\alpha}$; 第二个集合有最大元素, 设为 $\bar{\beta}$ 。严格偏好关系不具有反身性, 所以对任何 λ , $[\lambda p + (1-\lambda)r]Pq$ 和 $qP[\lambda p + (1-\lambda)r]$ 不可能同时成立。因此, $\bar{\beta} \leq \underline{\alpha}$ 。对于 $\lambda \in [\bar{\beta}, \underline{\alpha}]$, $[\lambda p + (1-\lambda)r]Pq$ 和 $qP[\lambda p + (1-\lambda)r]$ 都不成立。因此, 根据完全性公理, 对所有的 $\lambda \in [\bar{\beta}, \underline{\alpha}]$, $[\lambda p + (1-\lambda)r]Iq$ 。□

再来探讨独立性公理——这可能是四个公理中最有争议的一个。^①假设某个决策者在两个概率分布之间有某种偏好序。如果这两个分布分别(以相同的概率)与另一个分布结合; 根据独立性公理, 决策者的偏好序与两个原始概率分布上的偏好序相同。设想两个概率分布: 第一个分布以 0.5 的概率产生 100 元的收益, 以 0.5 的概率产生 0 元钱的收益; 第二个概率分布则产生 75 元的确定收益。我们在这两个概率分布的基础上分别构造如下两个复合分布: 以 0.5 的概率产生 100 万元的收益, 以 0.5 的概率产生原来的概率分布。根据独立性公理, 决策者在这两个复合概率分布上的偏好, 与在两个原始分布上的偏好相同。借用概率分布的树状图来解释, 独立性公理说的是, 决策者对两个概率分布的排序, 只决定于他对这两个分布中不同的结果分支的排序。这就是为什么在第一个例子中我们不考虑 x_2 。下面两个引理把独立性公理推广至无差异和严格偏好关系中。

引理 3.2 对任意的 $\alpha \in (0, 1)$ 和概率分布 $r, p, q \in P$, pIq 当且仅当 $[\alpha p + (1-\alpha)r]I[\alpha q + (1-\alpha)r]$ 。

证明: 我们证明充分性。设 pIq 。根据独立性公理, $[\alpha p + (1-\alpha)r]R[\alpha q + (1-\alpha)r]$ 且 $[\alpha q + (1-\alpha)r]R[\alpha p + (1-\alpha)r]$ 。因此, $[\alpha p + (1-\alpha)r]I[\alpha q + (1-\alpha)r]$ 。同理可证明必要性。□

引理 3.3 对任意的 $\alpha \in (0, 1)$ 和概率分布 $r, p, q \in P$, pPq 当且仅当 $[\alpha p + (1-\alpha)r]P[\alpha q + (1-\alpha)r]$ 。

证明: 我们证明充分性。设 pPq 。根据独立性公理, $[\alpha p + (1-\alpha)r]R[\alpha q + (1-\alpha)r]$ 。现在证明无差异性违背独立性公理。设 $[\alpha q + (1-\alpha)r]I[\alpha p + (1-\alpha)r]$ 。根据引理 3.2, qIp , 而这与前提条件 pPq 矛盾。类似方法可证明必要性。□

独立性公理还有一个重要的但并不如此直接的结论:

引理 3.4 如果 pRq 且 $\alpha \in (0, 1)$, 则 $pR[\alpha p + (1-\alpha)q]Rq$ 。

证明: 设 pRq 。复合概率分布的简化、引理 3.3 和偏好的传递性意味着:

$$pI[\alpha p + (1-\alpha)p]P[\alpha q + (1-\alpha)q]P[\alpha q + (1-\alpha)q]Iq$$

① 在一些教材和文章中, 独立性公理被称为替代公理。

如果 pIq , 根据引理 3.2, 通过轮换使用 $r = p$ 和 $r = q$, 得到 $pI[ap + (1-\alpha)q]Iq$. 因此, $pR[ap + (1-\alpha)q]Rq$. $\square$

根据这一引理, 对两个概率分布进行加权平均, 由此得到的概率分布在偏好排序中处于中间位置。独立性和连续性公理的另一个结论则对期望效用函数的存在起着至关重要的作用。

引理 3.5 设各个备选方案通过下标区别开来, 对所有的 j, x_1Rx_j 且 x_jRx_j . 则对所有的 $\alpha, \beta \in [0, 1]$,

$$[\alpha x_1 + (1-\alpha)x_j]R[\beta x_1 + (1-\beta)x_j] \text{ 当且仅当 } \alpha \geq \beta$$

证明: 设 $\alpha \geq \beta$. 我们可以把 $\alpha x_1 + (1-\alpha)x_j$ 写为 $\gamma x_1 + (1-\gamma)[\beta x_1 + (1-\beta)x_j]$, 其中, $\gamma = (\alpha - \beta)/(1 - \beta) \in (0, 1)$. 根据引理 3.4, $x_1R[\beta x_1 + (1-\beta)x_j]$. 再次应用引理 3.4, 得到 $[\gamma x_1 + (1-\gamma)[\beta x_1 + (1-\beta)x_j]]R[\beta x_1 + (1-\beta)x_j]$.

必要条件的证明与此类似, 只是 α 和 β 的位置要互换. $\square$

直观地看, 在比较最好的结果与最差的结果上的概率分布时, 决策者偏好能以最大的概率产生最好的结果的概率分布。上述公理和引理可用于证明概率分布上的偏好能用期望效用函数表示。我们用概率分布上的偏好将这一结论表述为下面的定理。由于是行动诱生结果上的概率分布, 所以我们还可以用行动上的偏好表述这一定理。

定理 3.1 (von Neumann-Morgenstern) 如果公理 3.1 到公理 3.4 都成立, 则存在函数 $u(x_j)$, 它给每个结果 x_j 赋予一个实数 u_j , 从而使得由行动 i 诱生的概率分布 p_i 的期望效用为:

$$EU(p_i) = p_{i1}u_1 + p_{i2}u_2 + \cdots + p_{ij}u_j = \sum_{j=1}^J p_{ij}u_j$$

并且, p_iRp_j 当且仅当 $EU(p_i) \geq EU(p_j)$ 。

函数 $u(x_j)$ 有时被称为贝努里(Bernoulli)效用函数, 借此把它与期望效用函数 $EU(p)$ 区别开来。贝努里函数和期望效用函数是两个完全不同的概念。前者定义在结果上, 后者定义在概率分布上。

定理 3.1 指出, 概率分布的期望效用是对结果带来的效用的加权平均, 权重是各个结果的发生概率。在前面的例子中, 如果赋予结果 x_1 、 x_2 和 x_3 的效用分别为 $u(x_1)$ 、 $u(x_2)$ 和 $u(x_3)$, 则概率分布 p 的期望效用为:

$$EU(p) = \left(\frac{1}{3}\right)u(x_1) + \left(\frac{1}{2}\right)u(x_2) + \left(\frac{1}{6}\right)u(x_3)$$

概率分布 q 的期望效用为:

$$EU(q) = \left(\frac{1}{4}\right)u(x_1) + \left(\frac{1}{2}\right)u(x_2) + \left(\frac{1}{4}\right)u(x_3)$$

期望效用函数具有的一个性质是, 它是贝努里效用的线性函数。根据这一性质, 如

果 r 是复合概率分布, 以 $\frac{1}{4}$ 的概率产生分布 p , 以 $\frac{3}{4}$ 的概率产生分布 q , 则:

$$EU(r) = EU\left(\frac{1}{4}p + \frac{3}{4}q\right) = \frac{1}{4}EU(p) + \frac{3}{4}EU(q) = \frac{13}{48}u(x_1) + \frac{1}{2}u(x_2) + \frac{11}{48}u(x_3)$$

这一结论和简化概率分布产生的期望效用完全相同。

这里不对 von Neumann-Morgenstern 定理给出数学证明, 而仅就三种结果的情况给出证明思路, 使用数学归纳法, 可以证明一般结论。作为习题, 请你证明一般结论。设 $X = \{x_1, x_2, x_3\}$, $x_1 R x_2 R x_3$, 且其中至少有一个严格偏好关系, 因为 $x_1 I x_2 I x_3$ 的情况很容易处理。我们把 X 上的概率分布表示为向量形式 (p_1, p_2, p_3) 。根据引理 3.1, 存在 α 使得 $x_2 I(\alpha, 0, 1-\alpha)$ 。同理, $x_1 I(1, 0, 0)$ 和 $x_3 I(0, 0, 1)$ 。设 $u_1 = 1, u_2 = \alpha$ 和 $u_3 = 0$ 。现在分析任意概率分布 $p = (p_1, p_2, p_3)$ 。从这一概率分布出发, 用概率分布 $(u_2, 0, 1-u_2)$ 替代使 x_2 肯定发生的退化概率分布, 得到复合概率分布 $(p_1 + u_2 p_2, 0, p_3 + p_2(1-u_2))$ 。根据引理 3.2, 决策者在这一复合概率分布和 p 之间肯定无差异。再对 x_1 和 x_3 进行替代。我们就能看到, 决策者在 p 和 $\{p_1 u_1 + p_2 u_2 + p_3 u_3, 0, p_1(1-u_1) + p_2(1-u_2) + p_3(1-u_3)\}$ 之间无差异。分析另一任意概率分布 q 。重复上述步骤。我们就能看到, 决策者在 q 和 $\{q_1 u_1 + q_2 u_2 + q_3 u_3, 0, q_1(1-u_1) + q_2(1-u_2) + q_3(1-u_3)\}$ 之间无差异。

根据引理 3.5, $p R q$ 当且仅当 $p_1 u_1 + p_2 u_2 + p_3 u_3 \geq q_1 u_1 + q_2 u_2 + q_3 u_3$ 。于是, 概率分布 p 的效用可以用 $p_1 u_1 + p_2 u_2 + p_3 u_3$ 表示。但是, 这一定理并没有说任意的 $\alpha \in [0, 1]$ 在这里都行得通, 而只是保证至少存在这样的 α , 使得结果效用 $u_1 = 1, u_2 = \alpha$ 和 $u_3 = 0$ 在这里能起作用。

基数效用

在前一章中, 我们告诉效用理论怀疑者, “放松些, 效用函数不过是对偏好序的表示而已”。但是, 对期望效用而言, 这种说法不再成立。结果上的效用 $u(x_j)$ 不再只是序数, 而且是基数, 因为它们给出了关于结果上的相对偏好的信息。就像华氏温度计告诉我们的, 212 度和 32 度之间的温差, 是 122 度和 32 度之间的温差的两倍一样, 基数效用函数使得我们能够说, “相对于鸡肉我对牛排的偏好程度, 是相对于鱼肉我对鸡肉的偏好程度的 3.8 倍”。与序数效用不同, $u(x_j) - u(x_i)$ 现在有了内涵。

期望效用理论为什么决定于基数贝努里效用函数呢? 假设有人要在结果 x_1, x_2 和 x_3 上的两个概率分布之间做出选择。假设此人在三个结果上的偏好序为 $x_1 P x_2 P x_3$ 。在概率分布 1 中, x_1 的发生概率是 0.5, x_3 的发生概率也是 0.5。在概率分布 2 中, 结果 x_2 是确定的结果。设 $u(x_1) = 1, u(x_2) = \alpha \in (0, 0.5), u(x_3) = 0$ 。这一效用表示告诉我们, 此人会选择概率分布 1。如果这一效用表示只具有序数特征, 我们就可以对效用函数进行任意转换, 只要它的序关系保持不变,

转换得到的函数依然表示原来的偏好。我们看对前述贝努里效用的一个转换： $v(x_1) = 1$, $v(x_2) = 1 - \alpha$, $v(x_3) = 0$ 。这一转换保留了结果上的偏好序,但是,此人现在偏好概率分布 2。

华氏温度计不是提供关于温度的相对信息的唯一方法;期望效用表示也并非只有唯一一个。设有两个基数效用函数 $u(x_j)$ 和 $v(x_j) = a + bu(x_j)$, 其中 $b > 0$ 。给定概率分布 p , 这两个基数效用函数分别产生期望效用 $\sum_{j=1}^J p_j u(x_j)$ 和 $a + b \sum_{j=1}^J p_j u(x_j)$ 。由于为 $\sum_{j=1}^J p_j u(x_j)$ 和 $\sum_{j=1}^J q_j u(x_j)$ 之间的序关系与 $a + b \sum_{j=1}^J p_j u(x_j)$ 和 $a + b \sum_{j=1}^J q_j u(x_j)$ 之间的序关系相同,所以这两个基数效用函数产生完全相同的行为。这一结论同时告诉我们, $u(x_j) - u(x_k)$ 在规模因子 b 下保持不变,且相对差 $\frac{u(x_j) - u(x_k)}{u(x_l) - u(x_m)}$ 唯一。所以说,贝努里效用函数不仅仅是序数的,而且服从正线性转换(或者说,仿射转换)。

3.2 风险偏好

在不确定性条件下的选择中,有一个性质无法用前一节中的公理描述,这就是理性决策者愿意承受的风险程度。一些人希望得到好的结果,即使好的结果发生的概率很小;还有一些人希望把差的结果的发生概率降低到最小,即使他得为此放弃获得高收益的机会。前面说过,给定三个结果 x 、 y 和 z , 如果 $xPyPz$, 根据连续性公理,与确定地得到 y 相比,决策者偏好 x 和 z 上的某一概率分布,如果 z 的发生概率足够小的话。但是,这一公理没有告诉我们 z 的发生概率需要小到什么程度。

分析决策者对风险的偏好的常用方法,是看他是否愿意接受公平赌局。所谓公平赌局,指的是其价格或赌注等于其期望收入的赌局。这里假设赌局的价格和收入(即结果)用货币表示,或者用决策者认为越多越好的某种商品表示。在后面,我们将探讨决策者有匮乏性偏好的情况——只在某个水平以下,决策者才希望越多越好。

设 w 是赌局的赌注, x_1 和 x_2 是赌局的两个货币结果, $x^1 > x^2$ 。设 p 是 x_1 的发生概率, $1 - p$ 是 x_2 的发生概率。我们有下面的定义:

定义 3.1 如果 $w = px_1 - (1 - p)x_2$, 则这个赌局是公平的;如果 $w < px_1 - (1 - p)x_2$, 则这个赌局是有利的;如果 $w > px_1 - (1 - p)x_2$, 则这个赌局是不公平的。

一张彩票,如果它以 $\frac{1}{100}$ 的概率产生 100 元钱的收入,以 $\frac{99}{100}$ 的概率产生 0 元钱的收入,它的价格为 1 元钱,它就是一个公平赌局。如果它的价格不到 1 元钱,它就是有利的赌局;如果它的价格高于 1 元钱,它就是不公平的赌局。在现实世界上,被称为“彩票”的几乎所有的彩票(特别是美国的州政府发行的彩票)都是不公平的。利用公平赌局这一概念,我们可以定义对风险的偏好。

定义 3.2 决策者如果不接受公平赌局,就是说,如果对所有的 p 、 x_1 和 x_2 , $u(px_1 - (1-p)x_2) > pu(x_1) - (1-p)u(x_2)$, 他就是厌恶风险的。

定义 3.3 决策者如果接受公平赌局,就是说,如果对所有的 p 、 x_1 和 x_2 , $u(px_1 - (1-p)x_2) < pu(x_1) - (1-p)u(x_2)$, 他就是喜好风险的。

定义 3.4 决策者如果在任意的公平赌局和这一赌局的赌注之间是无差异的,就是说,如果对所有的 p 、 x_1 和 x_2 , $u(px_1 - (1-p)x_2) = pu(x_1) - (1-p)u(x_2)$, 他就是风险中性的。

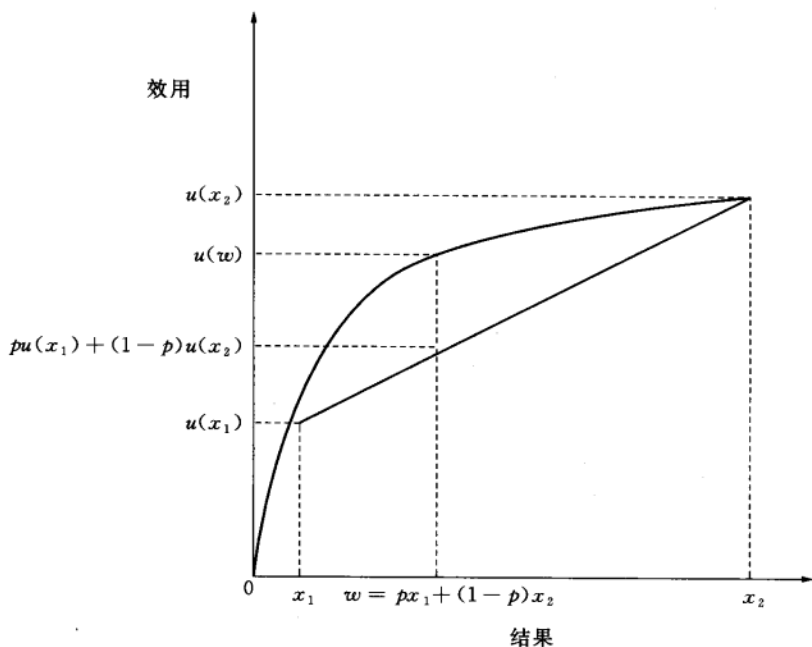


图 3.4 厌恶风险

从上述定义可以看出,决策者对风险的偏好与其效用函数的形状有密切的关系。图 3.4 给出了公平赌局 $w = px_1 - (1-p)x_2$ 的效用。连接点 $(x_1, u(x_1))$ 和 $(x_2, u(x_2))$ 的直线穿过点 $[pu(x_1) - (1-p)u(x_2), px_1 - (1-p)x_2]$ 。因此, $pu(x_1) - (1-p)u(x_2)$ 的值,位于 $(x_1, u(x_1))$ 和 $(x_2, u(x_2))$ 之间的直线与 w 点上的垂线的交点。我们能够看出, $u(w) > pu(x_1) - (1-p)u(x_2)$, 所以,这位决策者厌恶风险,不会接受这一公平赌局。效用函数之所以能产生这种结果,是因为它始终位于连接任意两个效用水平的直线的上方。函数的这一特征被称为凹性。

在图 3.5 中,效用函数始终位于连接任意两个效用水平的直线的下方。给定凸效用函数,如图 3.5 中的函数,决策者始终接受公平赌局,因为 $pu(x_1) - (1-p)u(x_2) > u(px_1 - (1-p)x_2)$ 。图 3.6 中的线性效用函数产生风险中性行为;给定这类效用函数,赌局的期望效用等于其期望结果产生的效用。

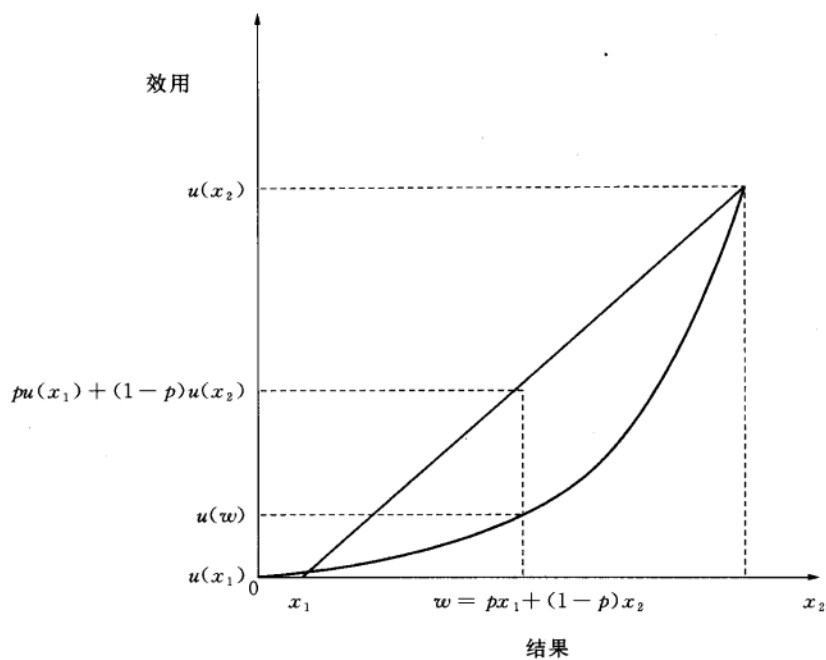


图 3.5 喜好风险

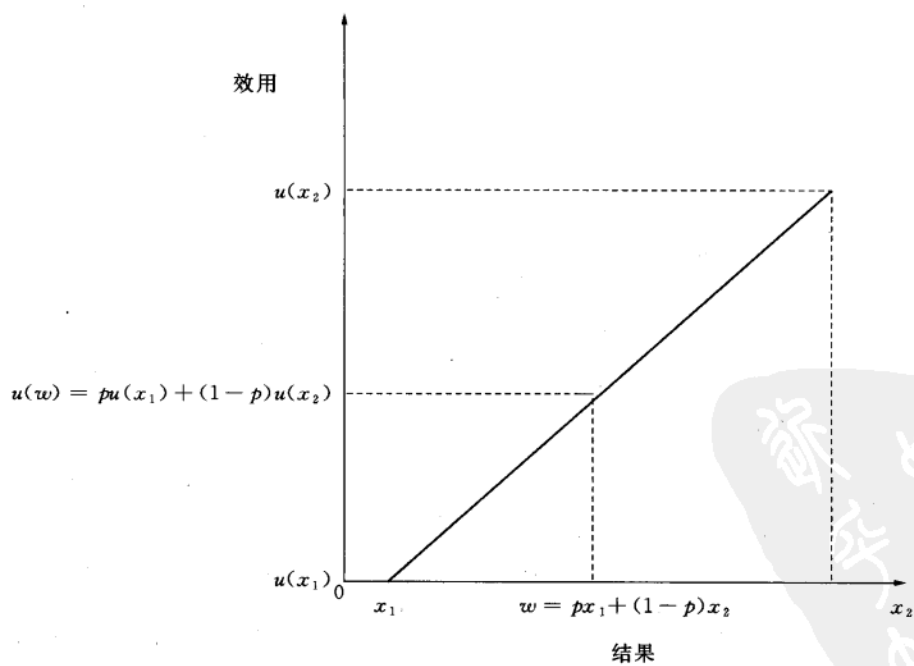


图 3.6 风险中性

3.2.1 风险偏好和随机占优

给定一个概率分布,只要它对任意数量的(有限个)可能的结果赋予正概率,我们就可以应用上述概念。

定义 3.5 如果对任意的非确定性的概率分布 p , $u(\sum_j p_j x_j) > \sum_j p_j u(x_j)$, 则偏好关系具有风险厌恶特征。

定义 3.6 如果对任意的非确定性的概率分布 p , $u(\sum_j p_j x_j) < \sum_j p_j u(x_j)$, 则偏好关系具有喜好风险特征。

要把风险厌恶与行为联系在一起,我们还需要更多的定义。给定一个概率分布 p , 它的期望值为 $E(p) = \sum_j p_j x_j$, 它的方差为 $V(p) = \sum_j p_j (E(p) - x_j)^2$ 。给定概率分布 p 和期望效用表示 $u = (u_1, \dots, u_j)$, 则 p 的确定性等价 $C(p)$ 是决策者认为和这一概率分布具有相同效用的结果。就是说,

$$u(C(p)) = EU(p)$$

决策者的风险厌恶行为在一定程度上是可以预测到的。他们愿意为减少方差而放弃部分期望值,因此,确定性等价小于期望值。如果两个概率分布 p 和 q 有相同的期望值,但是 q 的方差更大,则决策者偏好 p 。在这种情况下, q 是 p 的均值保留展型。

定义 3.7 如果 q 等于概率分布 p 加上某个随机项 ε , 则 q 就是个复合概率分布;如果 ε 有分布 z 并且 $E(z) = 0$ 且 $V(z) > 0$, 则概率分布 q 是 p 的均值保留展型(mean-preserving spread)。

定理 3.2 给定各个概率分布上的偏好关系 R 和表示 R 的贝努里效用函数 $u(x)$, 则以下各个陈述等价。

- (1) 偏好关系 R 具有风险厌恶特征。
- (2) 效用函数 $u(x)$ 是严格凹的。
- (3) 对任意概率分布 p , $C(p) \leq E(p)$ (如果 $V(p) > 0$, 这一不等式为严格不等式)。
- (4) 对任意两个概率分布 q 和 p , 如果 q 是 p 的均值保留展型, 则 $EU(q) < EU(p)$ 。

证明: 陈述(1)和陈述(2)的等价性可谓一目了然。

我们来证明陈述(2)意味着陈述(3)。设 R 具有风险厌恶特征。就是说, 对任何非确定性概率分布 p (例如 $V(p) > 0$), $u(E(p)) > \sum_j p_j u(x_j)$ 。但是, $u(C(p)) = EU(p)$ 意味着 $u(E(p)) > u(C(p))$ 。由于 $u(x)$ 是增函数, 所以, $E(p) > C(p)$ 。

我们再来证明陈述(3)意味着陈述(2)。假设陈述(3)成立。给定任意两种结果 x_1 和 x_2 , 设 x_1 的发生概率为 $p \in (0, 1)$ 。根据陈述(3)和 $u(x)$ 是增函数的事实,

$u(px_1 + (1-p)x_2) > pu(x_1) + (1-p)u(x_2)$, 即函数 $u(x)$ 是凹函数。在 $V(p) = 0$ 时(即 p 是确定性概率分布时), 这两个表达式相等。于是, 命题得证。

我们再来证明陈述(2)意味着陈述(4)。给定 p , 设 q 是均值保留展型, 它等于 $x + \varepsilon$, 其中, x 有分布 p , $\varepsilon \in \{\varepsilon_1, \dots, \varepsilon_T\}$ 有分布 z 。因此,

$$EU(q) = \sum_i \sum_j z_i p_j u(x_j + \varepsilon_i)$$

根据陈述(2), 对 x_j 的每个值, 有:

$$\sum_i z_i p_j u(x_j + \varepsilon_i) < u(\sum_i z_i p_j (x_j + \varepsilon_i))$$

调整不等号右边, 得到:

$$\sum_i z_i p_j u(x_j + \varepsilon_i) < u(p_j x_j + \sum_i z_i \varepsilon_i) = u(p_j x_j)$$

就 j 求和, 得到:

$$EU(q) = \sum_i \sum_j z_i p_j u(x_j + \varepsilon_i) < \sum_j u(p_j x_j) = EU(p)$$

我们再来证明陈述(4)意味着陈述(3)。分析概率分布 q 。 q 是概率分布 p 的均值保留展型, 它赋予 $E(q)$ 的概率为 1。命题 4 意味着 $EU(p) > EU(q)$ 。由于 p 是确定性概率分布, $u(E(p)) = EU(p)$, 于是 $u(E(p)) > EU(q)$ 。 $u(x)$ 的单调性和 $u(C(q)) = EU(q)$ 的事实, 意味着 $E(p) = E(q) > C(q)$ 。 $\square$

如果 q 是 p 的均值保留展型, 则相对于 q , p 是二阶随机占优。重复应用前面的证明思路, 我们还能得到下面的结论, 它简洁地概括了二阶随机占优的性质。

定理 3.3 相对于概率分布 q , p 是二阶随机占优的, 当且仅当对每个凹的增函数 $u(x)$, 有:

$$\sum_j p_j u(x_j) \geq \sum_j q_j u(x_j)$$

这就是说, 风险厌恶者的偏好等价于二阶随机占优——如果相对于 q , p 是二阶占优的, 则风险厌恶者偏好 p 。有时, 概率分布之间的选择十分简单。例如, 给定两个概率分布, 如果其中一个相对于另一个是一阶随机占优, 这时, 在两者之间的选择便“不用动脑子”。

定义 3.8 相对于概率分布 q , 概率分布 p 是一阶随机占优的, 如果对任意非递减函数 $u(x)$,

$$\sum_j p_j u(x_j) \geq \sum_j q_j u(x_j)$$

就是说, 如果概率分布之间存在着一阶随机占优关系, 则在它们之间做出选择, 与风险态度没有任何关系。

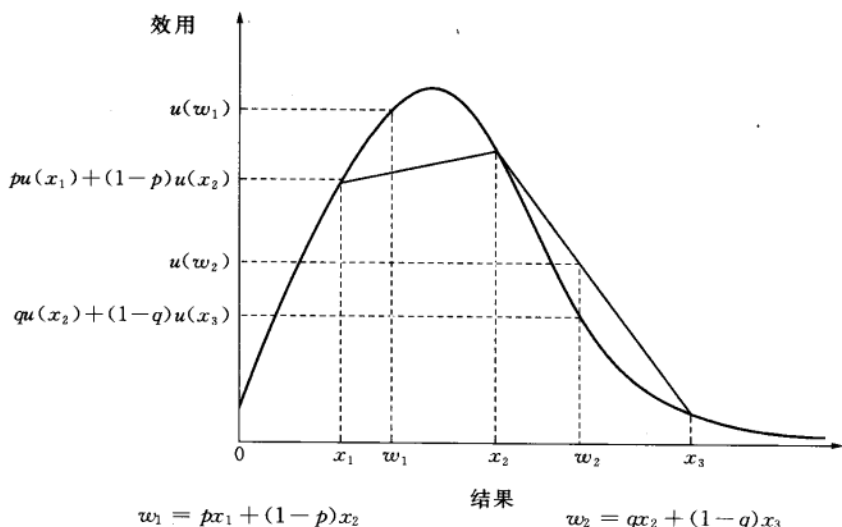


图 3.7 风险和空间偏好

3.2.2 有满足性偏好的风险偏好

第2章说过,政治科学中用到的许多效用函数具有满足性,因为人们有最偏好的结果。这种偏好所具有的风险承受能力,自然值得研究。图3.7给出了一个效用函数和结果 x_1 与 x_2 上的一个概率分布。其中, x_1 小于理想点, x_2 大于理想点。这个概率分布的期望效用是 $(x_1, u(x_1))$ 和 $(x_2, u(x_2))$ 之间的线段与 w 点的垂线的交点。由于理想点在 x_1 与 x_2 之间,所以在这个区间中肯定存在一个结果 w_1 ,使得有 $u(w_1) > pu(x_1) + (1-p)u(x_2)$ 。因此,在理想点附近满足性偏好具有厌恶风险特征。但是,它未必具有全局性风险厌恶特征。在远离理想点的区域中,如果效用函数是凸的,则这个区域上的概率分布具有喜好风险的特征。下一节将更详细地讨论这一点。

3.2.3 风险和高维空间中的欧几里得偏好*

在本节中,我们在 $\mathbb{R}^n$ 空间中讨论概率分布上的选择。Bendor 和 Meirowitz (2004)指出,我们可以把在严格递增偏好情况中的风险厌恶概念推广至欧几里得偏好。前面说过, $x \in \mathbb{R}^n$ 上的偏好是欧几里得偏好,如果我们能把这种偏好表示为如下形式的效用函数:

$$u(x) = h(-\|x - x^*\|)$$

其中, x^* 是 $\mathbb{R}^n$ 中的点, $h(\cdot)$ 是严格增函数。如果函数 $h(\cdot)$ 是严格凸的,则效用函数 $u(x)$ 是严格凹的。于是,贝努里效用函数反映风险厌恶偏好。即使函数 h 不

是凸的,偏好依然具有风险厌恶特征。为把均值保留展型概念推广至 $\mathbb{R}^n$,我们直接应用定义 3.7,只是此时的状态发生在 $\mathbb{R}^n$ 中。

定理 3.4 如果 $u(x)$ 表示欧几里得偏好,理想点为 x^* ,则对 $\mathbb{R}^n$ 上的任意两个概率分布 q 和 p ,如果它们有期望值 x^* 且 q 是 p 的均值保留展型,则 $EU(q) < EU(p)$ 。

证明: 设 q 和 p 是 $\mathbb{R}^n$ 上的概率分布,都有期望值 x^* ,且 q 是 p 的均值保留展型。定义 $d_j \equiv \|x_j - x^*\|$ 为结果 x_j 和点 x^* 之间的距离。由于偏好是欧几里得偏好,所以, $EU(q) < EU(p)$, 当且仅当存在某个递增函数 $h(\cdot)$,使得:

$$\sum_j q_j h(-d_j) < \sum_j p_j h(-d_j)$$

由于 q 是 p 的均值保留展型,所以:

$$\sum_j q_j d_j > \sum_j p_j d_j$$

于是,对任意的增函数 $g(\cdot)$,包括 $h(\cdot)$,

$$\sum_j q_j g(-d_j) < \sum_j p_j g(-d_j) \quad \square$$

这一定理指出,对以理想点为中心的概率分布,结果上的二阶随机占优等价于负效用上的一阶随机占优。或者说,有欧几里得偏好的决策者,在以其理想点为中心的概率分布上都厌恶风险。

3.3 学习

我们需要知道,在得到关于各个结果的发生概率的新信息时,理性决策者会如何做出反应。首先看一个例子。在图 3.8 中,一位投票者认为,在位的政治家是位“好”政治家的概率为 $\frac{3}{4}$,是位“差”政治家的概率是 $\frac{1}{4}$ 。假设这位投票者又得到了关于这位政治家任职表现的信息,例如其经济政策所导致的通货膨胀率。这一信息如何改变他对这位政客特征的概率推断呢?

假设此人知道,更好的政治家能以更大的概率实现更低的通货膨胀。具体地说,他知道,好的政治家能以 $\frac{2}{3}$ 的概率实现低通胀,差的政治家只能以 $\frac{1}{5}$ 的概率实现低通胀。直觉告诉我们,在通胀率低时,此人应该在初始概率 $\frac{3}{4}$ 的基础上,提高他对在位者是位好政治家的概率推断。在通胀率很高时,他应该降低对在位者是位好政治家的概率推断。我们还能把这种分析进行得更加深入,根据所实现的通胀率,计算在位者是“好”政治家的准确概率。

假设通货膨胀率很低。理性的投票者知道,这一结果出现在图 3.8 中第 2 幅图的最上面一个节点上或者出现在第三个节点上。而且他知道,到达最上面节点

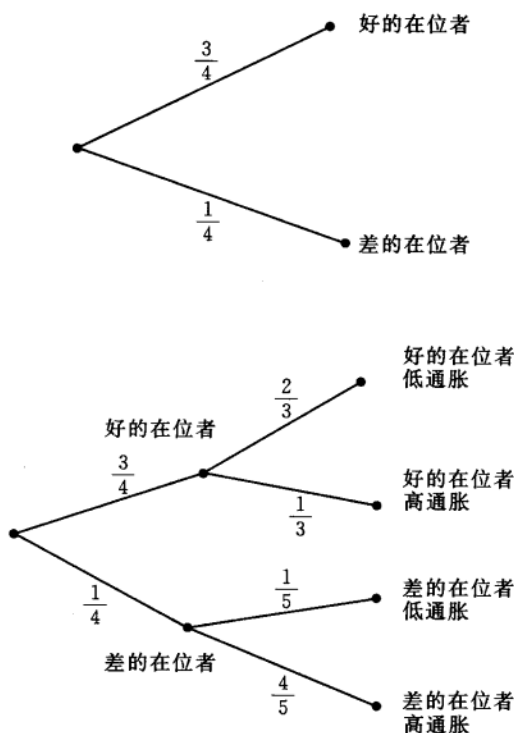


图 3.8 了解在位者的类型

的概率是 $\frac{3}{4} \times \frac{2}{3} = \frac{1}{2}$ ，到达第三个节点的概率是 $\frac{1}{4} \times \frac{1}{5} = \frac{1}{20}$ 。因此，在观察到低通胀后，投票者知道，在位者是位好政治家的可能性比他是位差政治家的可能性，高出 10 倍。设 $p(l)$ 是低通胀状态中在位者是位好政治家的概率。由于各结果的概率之和必须为 1，即 $p(l) + \frac{p(l)}{10} = 1$ ，因此， $p(l) = \frac{10}{11}$ 。同样的思路，能够得到 $p(h) = \frac{4}{9}$ 。这就证实了我们的直觉：低通胀率提高了在位者是位好政治家的概率，而高通胀率则降低了这一概率。

要把这个例子扩展成为一个学习模型，需要用到更精确的概率论概念。设 A 和 B 为两个事件（如图 3.8 中的终节点）。假设某个人观察到事件 B 发生，要据此计算事件 A 的发生概率——给定事件 B ，事件 A 的条件概率。我们将其定义为：

$$\Pr(A|B) = \frac{\Pr(A \& B)}{\Pr(B)}, \text{ 假设 } \Pr(B) > 0$$

其中， $\Pr(A)$ 是事件 A 的发生概率， $\Pr(B)$ 是事件 B 的发生概率， $\Pr(A \& B)$ 是事件 A 和 B 同时发生的概率（联合概率）。这一公式被称为贝叶斯法则，只在 $\Pr(B) \neq 0$ 时有定义。两个事件如果相互独立，则 $\Pr(A \& B) = \Pr(A)\Pr(B)$ ，

于是:

$$\Pr(A|B) = \frac{\Pr(A)\Pr(B)}{\Pr(B)} = \Pr(A)$$

在上面的例子中,我们应用贝叶斯法则进行分析。事件“低通胀和好的在位者”的发生概率 $\frac{1}{2}$,是给定“好的在位者”时低通胀的发生概率 $\frac{2}{3}$ 乘以“好的在位者”的发生概率 $\frac{3}{4}$ 。根据上述定义,我们有下面的定理:

定理 3.5 (贝叶斯法则) 设 $A_1 \cdots A_N$ 是不交事件(就是说,其中任意两个事件不同时发生)且对所有的 n , $\sum \Pr(A_n) = 1$ 且 $\Pr(A_n) > 0$ 。设 B 为任意一个事件。则:

$$\Pr(A_j|B) = \frac{\Pr(B|A_j)\Pr(A_j)}{\sum_{n=1}^N \Pr(B|A_n)\Pr(A_n)}$$

贝叶斯法则提供了一个简便易用的公式,我们能够用它计算在获得新信息后,理性决策者如何更新自己的概率推断。我们用它分析这里的投票者所面对的问题。事件 A_1 为在位者是位好政治家,事件 A_2 为在位者是位差政治家。由于在位者不可能同时既好且差,所以这两个事件不相交。事件 B 表示低通胀。于是,根据贝叶斯法则:

$$\Pr(A_1|B) = \frac{\Pr(B|A_1)\Pr(A_1)}{\Pr(B|A_1)\Pr(A_1) + \Pr(B|A_2)\Pr(A_2)}$$

$$\Pr(A_2|B) = \frac{\Pr(B|A_2)\Pr(A_2)}{\Pr(B|A_1)\Pr(A_1) + \Pr(B|A_2)\Pr(A_2)}$$

利用图 3.8 中的数字,我们得到以下概率:

$$\Pr(A_1) = \frac{3}{4}$$

$$\Pr(A_2) = \frac{1}{4}$$

$$\Pr(B|A_1) = \frac{2}{3}$$

$$\Pr(B|A_2) = \frac{1}{5}$$

把这些数字代入到上面两个条件概率表达式中,得到:

$$\Pr(A_1|B) = \frac{\frac{2}{3} \times \frac{3}{4}}{\frac{2}{3} \times \frac{3}{4} + \frac{1}{5} \times \frac{1}{4}} = \frac{10}{11}$$

$$\Pr(A_2|B) = \frac{\frac{1}{5} \times \frac{1}{4}}{\frac{2}{3} \times \frac{3}{4} + \frac{1}{5} \times \frac{1}{4}} = \frac{1}{11}$$

贝叶斯法则尽管有逻辑性,却常被认为是个很糟糕的学习模型。它不仅需要计算,并且超越了常人对条件概率的理解能力,而且其结果常与直觉相悖。我们来看由 Monte Hall 主持的“Let's Make a Deal”游戏中的一幕。Monte 让参赛者选择并打开三扇门中的一扇。在其中的一扇门后,藏着奖品,是一辆豪华轿车;在另两扇门后,只有很便宜的奖品(山羊似乎是这一节目中常用的奖品)。在参赛者做出选择之后和在打开所选择的这扇门之前,Monte 会打开另外两扇门中的一扇,从里面放出一只山羊。然后,他问参赛者是否想改变自己的选择,就是说,选择剩下的那扇紧闭的门。理性的参赛者应改变自己的选择吗?大部分人的直觉是,做这种改变是徒劳无功;改变选择而得到豪华汽车的概率,和在初始选择中得到这辆车的概率相同,赢得汽车的概率都是 $\frac{1}{3}$ 。这个问题曾公开在一份颇受欢迎的报纸专栏中,许多数学家和统计学家对此问题的看法与大部分人的直觉相一致。但是,这种认识与贝叶斯法则不一致。

假设参赛者选择了 3 号门。由于各扇门在事前是相同的,所以这里的分析适用于选择其他任何一扇门。首先看参赛者不改变自己的选择时赢得汽车的概率。它显然和汽车停在 3 号门后的初始概率相同,为 $\frac{1}{3}$ 。再看参赛者中途改变选择而赢得汽车的概率。设事件 A_1 、事件 A_2 和事件 A_3 分别表示汽车藏于 1 号、2 号和 3 号门后;事件 B_1 和事件 B_2 分别表示 Monte 打开 1 号门和 2 号门。由于 $\Pr(A_1) = \Pr(A_2) = \Pr(A_3) = \frac{1}{3}$,我们只需要对所有事件计算 $\Pr(B_i|A_j)$ 。Monte 从不打开藏有汽车的那扇门,所以 $\Pr(B_1|A_1) = \Pr(B_2|A_2) = 0$ 。假定在事件 A_3 中, Monte 随机打开藏有山羊的两扇门中的一扇。于是, $\Pr(B_1|A_2) = \Pr(B_2|A_1) = 1$ 且 $\Pr(B_1|A_3) = \Pr(B_2|A_3) = \frac{1}{2}$ 。

假设 Monte 打开了 2 号门。中途改变选择的参赛者赢得汽车的概率是:

$$\begin{aligned}\Pr(A_1|B_2) &= \frac{\Pr(B_2|A_1)\Pr(A_1)}{\Pr(B_2|A_1)\Pr(A_1) + \Pr(B_2|A_2)\Pr(A_2) + \Pr(B_2|A_3)\Pr(A_3)} \\ &= \frac{1 \times \frac{1}{3}}{1 \times \frac{1}{3} + 0 \times \frac{1}{3} + \frac{1}{2} \times \frac{1}{3}} = \frac{2}{3}\end{aligned}$$

如果 Monte 打开 1 号门,赢得汽车的概率同样是 $\Pr(A_2|B_1) = \frac{2}{3}$ 。所以,中途改变选择的参赛者以 $\frac{2}{3}$ 的概率赢得汽车,而坚持初始选择的参赛者赢得汽车的机

会只有 $\frac{1}{3}$ 。^①

那么,中途改变选择不会带来好处的直觉,为什么不正确呢?因为大部分人没有认识到 Monte Hall 从来不会暴露这辆汽车所在位置的事实所具有的意义。他没有打开某扇门的事实,被中途改变选择的参赛者在决策中利用,而坚持初始选择者却没有使用这一信息。

Monte Hall 问题确实揭示了贝叶斯学习模型存在的问题,但是,我们不能就此否定贝叶斯法则。因为,贝叶斯法则正确地指出,中途改变选择的参赛者赢得汽车的机会是 $\frac{2}{3}$ 。经常观看这档节目的人,即使没有计算条件概率,也会认识到参赛者应该改变自己的初始选择。所以,要支持贝叶斯法则,我们不妨这样理解:人们选择自己的行动,就好像他们已经进行了条件概率计算一样,即使他们只是遵循着从经验中获得的法则。

3.4 对期望效用理论的批判

本书用到的大部分模型是以期望效用理论为基础,但是,有很大一块有影响力的文献对期望效用理论持批判态度。但是,这些批判和它们所提出的取代期望效用理论的模型在政治博弈论中的应用,尚处于起步阶段。^②

3.4.1 风险、不确定性和主观概率

对期望效用理论的基本批判之一,建立在对风险和不确定性的区别上。经济学家 Frank Knight(1921)最早指出,期望效用理论所分析的是风险而不是不确定性。按他的解释,不确定性指的是人们缺乏足够的统计信息去形成对各种结果的发生概率的估计。换句话说,不确定性指的是人们不知道概率分布集 $\mathbf{P}$ 。统计学家 Leonard Savage(1954)对这种批判进行了反驳,他认为,人们有关于 $\mathbf{P}$ 的主观推断,这种推断可用于形成结果上的(主观)概率分布。

但是,实验证据却对不确定性能否简化为关于推断的推断提出了质疑。我们看最早由 Daniel Ellsberg(1961)提出的悖论。两个箱子都装有红色球和黑色球。在 1 号箱子中,装有 100 个红色球和黑色球,但是红色球的比例未知。在 2 号箱子中,装有 50 个红色球和 50 个黑色球。

① 这里给出的 Monte Hall 问题的解有些复杂,原因是我们要借此演示贝叶斯法则。还有更容易的证明:中途改变自己的选择的参赛者,只有在自己最初选择了正确的那扇门时,才会失败。就是说,改变选择的参赛者,得不到汽车的概率是 $\frac{1}{3}$,而赢得汽车的概率是 $\frac{2}{3}$ 。

② (与以期望效用理论为基础的模型相对立的)“行为模型”,近些年里在经济学中的应用愈来愈多(Camerer, 2003)。

实验对象只要捡出一个红色球就能得到 100 元。大部分人选择从 2 号箱子中摸球。但是,当实验方案改为捡出一个黑色球就奖励 100 元时,大部分人还是选择从 2 号箱子中摸球。在两个方案下人们都从 2 号箱子中摸球,违反了期望效用理论中的公理。根据期望效用理论,选择从 2 号箱子中摸红色球,表明实验对象认为 1 号箱子中红色球数量少于 50 个;选择从 2 号箱子中摸黑色球,表明实验对象认为 1 号箱子中黑色球数量少于 50 个。这些推断显然与事先就知道 1 号箱子中装有 100 个球的事实不一致。在两种实验方案下都从 1 号箱子中摸球的做法,同样违反期望效用理论。

3.4.2 Allais 悖论

期望效用理论的其他结论也在实验环境中被检验。这些实验研究表明,存在大量与期望效用理论不一致的决策行为。其中最早的也是得到最多研究的一个悖论,由法国经济学家 Maurice Allais 第一个提出。他通过实验发现,实验对象做出与独立性公理不一致的选择。

在实验中,实验对象首先被要求在概率分布 **a** 和概率分布 **b** 中做出选择,其中,

概率分布 **a**: 0.33 的概率赢得 2 500 元, 0.66 的概率赢得 2 400 元, 0.01 的概率得到 0 元钱。

概率分布 **b**: 确定性地得到 2 400 元。

面对着这两个选项,绝大多数人选择概率分布 **b**。Kahneman 和 Tversky (1979) 发现,面对着这样两种分布时,82% 的人选择概率分布 **b**。

实验对象接着被要求在概率分布 **c** 和概率分布 **d** 中做出选择,其中,

概率分布 **c**: 0.33 的概率赢得 2 500 元, 0.67 的概率赢得 0 元。

概率分布 **d**: 0.34 的概率赢得 2 400 元, 0.66 的概率赢得 0 元。

人们一般选择 **c**。Kahneman 和 Tversky 发现,83% 的人选择这一概率分布。在第一个实验中选择 **b** 和在第二个实验中选择 **c**, 违背了独立性公理, 因此违背了期望效用理论。选择 **b** 和 **c** 意味着:

$$u(2\,400) > 0.33u(2\,500) + 0.66u(2\,400) + 0.01u(0)$$

$$0.33u(2\,500) + 0.67u(0) > 0.34u(2\,400) + 0.66u(0)$$

调整第一个不等式, 得到 $0.34u(2\,400) > 0.33u(2\,500) + 0.01u(0)$; 调整第二个不等式, 得到 $0.33u(2\,500) + 0.01u(0) > 0.34u(2\,400)$ 。我们得到了相矛盾的结果。这一矛盾违背了独立性公理。概率分布 **a** 是个复合概率分布——在结果 (2 500, 2 400, 0) 上, 概率分布为 $0.34\left(\frac{33}{34}, 0, \frac{1}{34}\right) + 0.66(0, 1, 0)$; 而概率分布 **b** 是

$0.34(0, 1, 0) + 0.66(0, 1, 0)$ 。如果 aPb , 根据独立性公理, $\left(\frac{33}{34}, 0, \frac{1}{34}\right)P(0, 1, 0)$; 而这意味着 $\left[0.34\left(\frac{33}{34}, 0, \frac{1}{34}\right) + 0.66(0, 1, 0)\right]P[0.34(0, 1, 0) + 0.66(0, 1, 0)]$, 就是说, cPd 。

3.4.3 期望(Prospect)理论^①

在其经典论文中, Kahneman 和 Tversky (1979) 提出了另一种决策模型, 解释 Allais 悖论和其他实验结论。此前, 许多人把 Allais 悖论归因于对确定性的偏好。Kahneman 和 Tversky 指出, 当所有的概率分布都远远偏离确定性时, 独立性公理也常被违反。他们设计了下面四个分布:

概率分布 **a**: 0.45 的概率得到 6 000 元, 0.55 的概率得到 0 元钱;

概率分布 **b**: 0.90 的概率得到 3 000 元, 0.10 的概率得到 0 元钱;

概率分布 **c**: 0.001 的概率得到 6 000 元, 0.999 的概率得到 0 元钱;

概率分布 **d**: 0.002 的概率得到 3 000 元, 0.998 的概率得到 0 元钱。

他们发现, 在 **a** 和 **b** 之间, 典型的选择是 **b**; 在 **c** 和 **d** 之间, 典型的选择是 **c**。这种选择违背了独立性公理。在概率分布 **c** 和概率分布 **d** 中, 高收益结果的发生概率都很小, 实验对象一般选择奖金数更高的那个概率分布。但是当两个分布中高收益的发生概率都比较高时, 实验对象一般选择相对更确定的那个概率分布。

但是, Kahneman 和 Tversky 指出, 在面对着的是损失上的概率分布而不是收益上的概率分布时, 对确定性的这种偏好并不成立。我们看下面几个概率分布:

概率分布 **a**: 0.80 的概率得到 -4 000 元, 0.20 的概率得到 0 元钱;

概率分布 **b**: 确定性地得到 -3 000 元;

概率分布 **c**: 0.20 的概率得到 -4 000 元, 0.80 的概率得到 0 元钱;

概率分布 **d**: 0.25 的概率得到 -3 000 元, 0.75 的概率得到 0 元钱。

如果 Allais 的悖论只是源于对确定性的偏好, **b** 和 **c** 就应该是典型选择。但是, Kahneman 和 Tversky 发现, 典型选择是 **a** 和 **d**。他们的解释是, 人们在收益上厌恶风险, 但是在损失上, 愿意承担风险。

Kahneman 和 Tversky 还指出, 概率分布的表述形式影响人们的选择。假设给某个人 1 000 元钱, 然后提供两个概率分布供他选择:

概率分布 **a**: 0.5 的概率得到 1 000 元, 0.5 的概率得到 0 元钱;

概率分布 **b**: 确定性地得到 500 元。

另一个人得到了 2 000 元, 然后被要求在下面两个概率分布之间做出选择:

^① Prospect Theory 在中文中有译为“展望理论”和“前景理论”, 此处译为“期望理论”, 但读者应将其与“期望效用理论”区别开。——译者注

概率分布 **c**: 0.5 的概率失去 1 000 元, 0.5 的概率得到 0 元钱;

概率分布 **d**: 确定性地失去 500 元。

Kahneman 和 Tversky 发现, **b** 和 **c** 是典型选择。

为解释这些发现, Kahneman 和 Tversky 提出了期望(Prospect)理论, 取代期望效用理论。在他们的模型中, 选择过程包含两个阶段: 编辑阶段和估价阶段。在编辑阶段中, 人们“组织和理解所面对的选项, 借此简化随后的估价和选择”。

3.4.3.1 编辑阶段

Kahneman 和 Tversky 认为在编辑阶段, 人们进行六项活动。

(1) 编码: Kahneman 和 Tversky 认为, 人们把收益和损失分开进行评估, 所以在编辑阶段, 第一步工作是确定参照点, 并据此把结果编码为收益或损失。

(2) 合并: 人们合并相同结果的发生概率。

(3) 分离: 人们找到并分离各个选项中的无风险成分。例如, 以 0.7 的概率产生 200 元和 0.3 的概率产生 100 元的概率分布, 被理解为 100 元的无风险收益和在另外 100 元的收入上的概率分布。

(4) 抵消: 在比较两个概率分布时, 人们不考虑它们之间的相同部分。例如, 在最后一个例子中, 2 000 元的奖金被“抵消”, 并不影响在 **c** 和 **d** 之间的选择。

(5) 简化: 人们可能会通过对概率四舍五入来简化计算, 例如把 0.49 记为 0.5, 或者不考虑极不可能发生的结果等。

(6) 找出占优方案: 人们不考虑一阶随机劣的概率分布。

3.4.3.2 估价阶段

Kahneman 和 Tversky 的估价模型与期望效用理论十分相似, 都假定人们用各个结果所对应的收益的加权平均数, 对概率分布进行估价。但是, Kahneman 和 Tversky 所使用的权重不是各个结果的主观概率而是这些主观概率的函数。与期望效用理论不同, 他们认为, 结果值函数应该不对称地处理收益和损失。

设 x 和 y 是两个不同的货币结果, p 是 x 的概率, q 是 y 的概率, 没有任何结果发生或者说收益为 0 的概率是 $1 - p - q$ 。如果 $x, y > 0$ 且 $p + q = 1$, Kahneman 和 Tversky 定义“期望”为严格正的; 如果 $x, y < 0$ 且 $p + q = 1$, 则为严格负的; 在其他所有情况中, 为正常的。对于正常的“期望”, 人们最大化

$$V(x, p; y, q) = \pi(p)v(x) + \pi(q)v(y)$$

其中, $v(x)$ 和 $v(y)$ 是每种结果的价值, $\pi(p)$ 和 $\pi(q)$ 是以结果的发生概率为基础的权重。他们假定 $v(0) = 0$, $\pi(0) = 0$, $\pi(1) = 1$ 。如果 v 是贝努里函数且对所有的 p , $\pi(p) = p$, 则函数 V 等价于期望效用函数。

对严格正的或严格负的“期望”, 例如 $x > y > 0$ 或 $x < y < 0$, 且 $p + q = 1$, 人们最大化

$$V(x, p; y, q) = v(y) + \pi(p)[v(x) - v(y)]$$

这一函数形式所反映的是,人们把概率分布分解为无风险部分 $v(y)$ 和风险部分 $v(x) - v(y)$ 。

期望理论的关键假设是,在收益和损失上, $v(\cdot)$ 是不对称的。Kahneman 和 Tversky 给出了三个具体假设:

- (1) 价值函数是用对参照点(即没有收益或损失)的偏离定义的;
- (2) 价值函数在收益上是凹函数,在损失上是凸函数;
- (3) 价值函数在损失上要比在收益上更加陡峭。

图 3.9 中的函数满足这些特征。

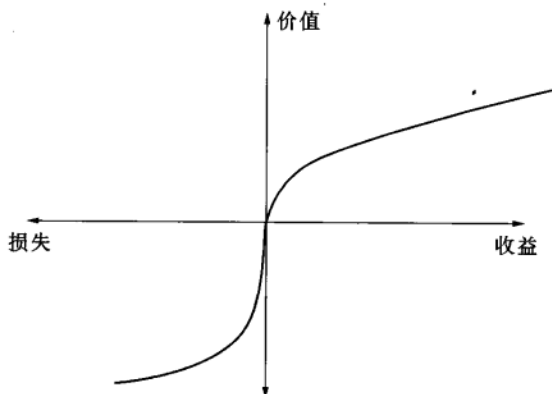


图 3.9 Prospect 理论:价值函数

另外,就决策权数 $\pi(p)$ 的形式, Kahneman 和 Tversky 给出了几个假设:

- (1) π 是 p 的增函数;
- (2) $\pi(0) = 0$;
- (3) $\pi(1) = 1$;
- (4) 在 p 的值较小时, $\pi(p) > p$;
- (5) 在 p 的值较小时, π 是次可加的:就是说, $\pi(rp) > r\pi(p)$, 其中, $0 < r < 1$;
- (6) 对所有的 p , π 具有确定性不足的特征:就是说, $\pi(p) + \pi(1-p) < 1$;
- (7) 对所有的 $0 < p, q, r < 1$, π 具有比例不足的特征:就是说, $\pi(pq)/\pi(p) \leq \pi(pqr)/\pi(pr)$ 。

前三个假设顺理成章。第四个假设指出,人们赋予小概率更大的权重。次可加性假设有助于解决 Allais 悖论;当它与假设 2 相结合时,意味着在比较小的 p 值上, π 是凸的(见习题 1 和习题 2)。确定性不足的假定也有助于解决 Allais 悖论。前面说过,典型的选择意味着 $v(2400) > \pi(0.33)v(2500) + \pi(0.66)v(2400)$ 和 $\pi(0.33)v(2500) > \pi(0.34)v(2400)$ 。这两个不等式意味着 $1 > \pi(0.66) + \pi(0.34)$ 。比例性不足的假设可以解释许多违背独立性公理的情况,因为对固定的

概率比例,这一假设意味着在概率很高时,决策权数的比重接近 1。

图 3.10 中画出了满足这些假设的函数。

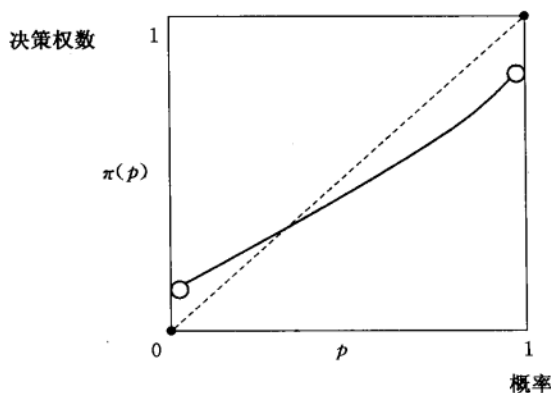


图 3.10 决策权数

3.5 时间偏好

本书中的许多动态模型要求决策者在当前收入和未来收入之间权衡取舍。我们可以合理地假定,相对于未来效用,人们更看重当前效用(在明天我们可能死去)。我们用贴现因子反映这一认识。设 $0 < \delta < 1$ 表示相对于当期效用,人们赋予下一时期效用的权重。这种常数贴现因子,意味着未来第二个时期的效用被赋予权重 δ^2 ,未来第 t 个时期的效用被赋予权重 δ^t ,如此等等。

3.5.1 计算收益流

我们常常把没有确定终点的博弈,描述为无限博弈,它的时期数趋于 ∞ 。在无限博弈中,我们不能简单地把各个时期中的收益相加,得到某个行动序列产生的效用。因为,这种固定效用流量的和为无穷大。幸运的是,几何贴现能方便我们处理这种情况。

我们看最简单的情况:在无限个时期中,某个人在每个时期中的效用为 $u_t = u$ 。我们有两种方法计算这个效用流量的值 v^∞ 。

方法 1:我们可以把 v^∞ 表示为 $u + \delta u + \delta^2 u + \cdots = u \sum_{t=0}^{\infty} \delta^t$ 。由于 $0 < \delta < 1$, $\sum_{t=0}^{\infty} \delta^t$ 是收敛幂级数。在 $t \rightarrow \infty$ 时, $\sum_{t=0}^{\infty} \delta^t$ 收敛于 $\frac{1}{1-\delta}$, 因此 $v^\infty = \frac{u}{1-\delta}$ 。对这个幂级数,还有以下结论:

$$\sum_{t=0}^T \delta^t = \frac{1 - \delta^{T+1}}{1 - \delta}$$

$$\sum_{t=T}^{\infty} \delta^t = \frac{\delta^T}{1-\delta}$$

$$\sum_{t=T}^S \delta^t = \frac{\delta^T - \delta^{S+1}}{1-\delta}$$

我们可以用上述结果计算有限效用流量。例如, T 个时期的效用流量为:

$$u \sum_{t=0}^T \delta^t = \frac{u(1-\delta^{T+1})}{1-\delta}$$

方法 2: 计算 v^∞ 的另一种方法是递归法, 它是 Bellman(1957) 最优性原理的基础。由于 v^∞ 是个无限期效用流量, 我们可以将其表示为单时期效用 u 加上从下个时期开始的无限期效用流量的贴现值, 即:

$$v^\infty = u + \delta v^\infty$$

因此, $v^\infty = \frac{u}{1-\delta}$ 。我们用这种方法计算有限期收益流量。假设在 T 个时期中的每个时期里, 决策者得到效用 u , 我们来计算 v^T 。由于 $v^T = v^\infty - \delta^{T+1} v^\infty$, 所以 $v^T = \frac{u}{1-\delta} - \frac{\delta^{T+1} u}{1-\delta} = \frac{u(1-\delta^{T+1})}{1-\delta}$ 。

在这个简单的例子中, 这种方法的优势没有显现, 但是, 在更复杂的环境中, 这种方法有显著优势。假设有 n 种自然状态 ($s_1, \dots, s_n$)。在状态 j , 决策者得到收益 u_j 。假设状态变化服从马尔科夫过程, 就是说, 在时期 t 中某状态的条件概率仅决定于时期 $t-1$ 中的状态。换句话说, 对每个 i, j 和 t , $\Pr(S_t = s_i | S_{t-1} = s_j) = \pi_{ij}$ 。

我们计算始自状态 j 的效用流量的值 v_j 。根据递归法:

$$v_j = u_j + \delta \sum_{i=1}^n \pi_{ij} v_i$$

这便有了由 n 个方程和 n 个未知数 (v_i) 构成的线性方程组。由于对所有的 j , $\sum_{i=1}^n \pi_{ij} = 1$, 所以, 用这一关系代替某个方程, 有时候更易求解出这一方程组。

我们看个简单的例子。假设我们想计算美国某个政党的长期收益, 它在执政时, 每年收益为 u_1 ; 在不执政时, 每年收益为 u_2 。假设存在着执政优势, 就是说, 当它执政时, 它在大选中胜出并继续执政 (即状态 1) 的概率为 $p > \frac{1}{2}$ 。这意味着如果它是在野党 (状态 2), 则在下个时期里, 它依然为在野党的概率是 p 。它从执政党变为在野党或者从在野党变为执政党的概率为 $1-p$ 。要计算这个政党在每种状态中的收益, 我们建立有关递归方程。我们有 $n=2$ 个未知数, $\pi_{11} = \pi_{22} = p$ 和 $\pi_{12} = \pi_{21} = 1-p$ 定义了方程组。假设 $u_1 > u_2$ 。于是, 两个方程为:

$$v_1 = u_1 + \delta(pv_1 + (1-p)v_2)$$

$$v_2 = u_2 + \delta((1-p)v_1 + pv_2)$$

求解这两个方程,得到:

$$v_1 = \frac{(1-\delta p)u_1 + \delta(1-p)u_2}{1-2\delta p + \delta^2(2p-1)}$$

$$v_2 = \frac{\delta(1-\delta p)u_1 + (1-\delta p)u_2}{1-2\delta p + \delta^2(2p-1)}$$

在上述例子中,效用流量是外生的——或者由常数效用构成或者由给定的概率分布生成。这一条件可以被放松。假设决策者从由状态决定的选择集 $X(s_t)$ 中选择政策 x_t , 最大化收益流 $u(x_t, s_t)$ 的贴现值, 其中, $s_t \in \{s_1, \dots, s_t\}$ 是时期 t 中的自然状态。我们假定自 s_t 开始的转换概率分布决定于 x_t , 于是, 在状态 s_t 和选择 x_t 之后, 出现状态 s_{t+1} 的概率为 $\pi(s_{t+1} | x_t, s_t)$ 。这里只探讨稳态方案(在这种方案中, 选择只决定于状态)。设 $x(s)$ 是稳态方案, 它规定了在状态为 s 时所采取的行动。

在状态 s 中执行方案 $x(s)$ 的收益为:

$$v(x(s), s) = u(x(s), s) + \delta \sum_{s'} v(x(s'), s) \pi(s' | x(s), s)$$

假设对所有方案, 我们求解出 $v(x(s), s)$, 于是最优收益为:

$$v^*(s) = \sup_x v(x(s), s)$$

根据 Bellman 的最优性原理,

$$v^*(s) = \sup_{x \in X(s)} [u(x, s) + \delta \sum_{s'} v^*(s') \pi(s' | x, s)]$$

3.5.2 双曲线贴现

大部分博弈论模型使用几何贴现, 但是在行为决策理论中, 越来越多的文献使用与试验证据更一致的另一种贴现方法。^① 这种广泛使用的方法便是双曲线贴现, 它假定在时点 0, 人们对时点 t 的效用的贴现因子为:

$$h(t) = (1 + \alpha t)^{-\frac{\gamma}{\alpha}}$$

其中, $\gamma > 0$, $\alpha > 0$ 。除非 α 接近 0, 否则双曲线贴现赋予未来效用的权重要比常数贴现大得多。它同时意味着人们有“时间一致性”问题。时点 t 的最优方案决定于时点 t 到现在的距离。假设某个人要决定如何把 1 元钱的消费支出分配到时期 0、1 和 2 中。设 $U(x) = \sqrt{x}$ 。在使用常数贴现因子时, 最优方案是下面的优化问题的解:

$$\max \sqrt{x_0} + \delta \sqrt{x_1} + \delta^2 \sqrt{x_2} \quad \text{s.t.} \quad \sum x_t = 1$$

最优解满足 $x_1 = \delta^2 x_0$ 和 $x_2 = \delta^4 x_0$ 。将这一条件代入到预算约束中, 得到:

① 在学习本节前, 你应该重温数学附录中的优化部分。

$$x_0 = \frac{1}{1 + \delta^2 + \delta^4}$$

$$x_1 = \frac{\delta^2}{1 + \delta^2 + \delta^4}$$

$$x_2 = \frac{\delta^4}{1 + \delta^2 + \delta^4}$$

假设在第一个时期之后,即在消费了 $x_0 = \frac{1}{1 + \delta^2 + \delta^4}$ 之后,此人重新优化自己的消费结构。他再次最优地选择了 $x_2 = \delta^2 x_1$ 。将这一条件代入到预算约束 $x_1 + x_2 = \frac{\delta^2 + \delta^4}{1 + \delta^2 + \delta^4}$, 得到:

$$x_1 = \frac{1}{1 + \delta^2} \times \frac{\delta^2 + \delta^4}{1 + \delta^2 + \delta^4} = \frac{\delta^2}{1 + \delta^2 + \delta^4}$$

$$x_2 = \frac{\delta^4}{1 + \delta^2 + \delta^4}$$

因此,这个人会继续执行原来的最优消费方案,他在时期 2 中的消费数量和原方案中的数量一样。

假设此人现在使用双曲线贴现,我们来分析他的配置问题。为使计算简单,设 $\alpha = \gamma = 1$ 。此人现在所求解的优化问题为:

$$\max \sqrt{x_0} + \frac{1}{2}\sqrt{x_1} + \frac{1}{3}\sqrt{x_2} \quad \text{s. t.} \quad \sum x_t = 1$$

最优解满足一阶条件: $x_1 = \frac{1}{4}x_0$, $x_2 = \frac{1}{9}x_0$ 。于是,最优解为:

$$x_0 = \frac{36}{49}$$

$$x_1 = \frac{9}{49}$$

$$x_2 = \frac{4}{49}$$

假设在消费了 x_0 之后,此人重新优化自己的消费结构。新的一阶条件为 $x_2 = \frac{1}{4}x_1$ 。代入预算约束 $x_1 + x_2 = \frac{13}{49}$ 中,得到:

$$x_1 = \frac{4}{5} \times \frac{13}{49} = \frac{52}{245} > \frac{9}{49}$$

$$x_2 = \frac{13}{245} < \frac{4}{49}$$

就是说,此人改变了自己的最优方案,把更多的消费转移到时期 1。造成这种变化的原因是,和时期 0 的消费决策相比,在时期 1 的消费决策中,时期 1 相对于时期 2 的权重更高。

双曲线贴现现在解释试验的异常结果和消费的时间特征等方面(例如,退休者的消费小于常数贴现模型所预测的消费数量)尽管十分有用,但在政治科学中,却很少被用到。^①

3.6 习题

习题 3.1 斯密是位众议员,她要决定是否竞选参议员。她认为自己获得党内提名的概率是 50%;如果获得提名,赢得参议员席位的概率为 40%。假设她从参议员席位中得到的效用为 W 。如果竞选失败,她只能打道回府,经营自己家里的二手车交易业务,得到效用 L 。如果不参加竞选、依然保住自己的众议员席位带给她的效用为 H 。

- (1) 利用概率树,描述竞选决策所涉及的概率分布。
- (2) 计算竞选参议员实现的期望效用。
- (3) H 必须比 W 和 L 低多少才能使斯密决定竞选参议员?

习题 3.2 证明定理 3.1。

习题 3.3 计算如下概率分布的期望收益:在连续五个时期中的每个时期里,某个人抛一枚硬币,如果硬币正面向上,他得到 1 元钱。换句话说,如果他连续 x 次使得硬币正面向上,他得到 x 元钱。

习题 3.4 在 Monte Hall 主持的节目中,假设他不是始终打开藏有山羊的门,而是随机选择打开哪一扇门,就是说,他可能会打开藏有轿车的那扇门。在他打开了藏有轿车的那扇门时,参赛者显然应该改变自己的选择,即选择这扇被打开的门;如果在他打开门后,出现的是山羊,参赛者还应该改变自己的选择,即选择另一扇关闭的门吗?

习题 3.5 某个国家卷入一场战争。在每个时期里,它都得进行一场战役,耗费成本 $f > 0$ 。这个国家赢得每场战役的概率是 π 。如果它连续赢得两场战役,它就永久地赢得这场战争,得到收益 $w > 0$ 。如果它连续输掉两场战役,它就永久地输掉这场战争,损失 $l = 0$ 。这个国家对未来时期的贴现因子是 δ 。

对应着连续赢得战役或连续输掉战役,有五种可能的状态。其中两种状态是终点状态,分别为赢得战争或输掉战争。对于每个非终点状态,计算继续这场战争的期望效用。 f 必须满足什么条件才能使这个国家不挑起这场战争? f 又必须满

^① 一个概念上的障碍是,在这种贴现方法中无限时域上的效用可能没有定义。假设要计算由固定效用 u 构成的无限期效用流量,级数 $\sum_{t=0}^{\infty} \delta^t u$ 必须收敛。但是,在参数 α 和 γ 的很多取值上,它并不收敛。

是什么条件才能使这个国家在输掉一场战役后便举手投降呢?

习题 3.6 证明 $\pi(p) > p$ 和次可加性条件意味着在比较小的 p 值上, 决策权重函数 π 是凸函数。

习题 3.7 分析下面的概率分布:

概率分布 **a**: 0.45 的概率得到 6 000 元, 0.55 的概率得到 0 元;

概率分布 **b**: 0.90 的概率得到 3 000 元, 0.10 的概率得到 0 元;

概率分布 **c**: 0.001 的概率得到 6 000 元, 0.999 的概率得到 0 元;

概率分布 **d**: 0.002 的概率得到 3 000 元, 0.998 的概率得到 0 元。

Prospect 理论所预测的选择是哪一项?

习题 3.8 Kahneman 和 Tversky 发现, **b** 和 **c** 是下面的实验中的典型选择:

实验 1: 给某个人 1 000 元钱, 然后让他在下面两个概率分布之间做出选择:

概率分布 **a**: 0.5 的概率再得到 1 000 元, 0.5 的概率得到 0 元;

概率分布 **b**: 确定性地得到 500 元。

实验 2: 给某个人 2 000 元钱, 然后让他在下面两个概率分布之间做出选择:

概率分布 **c**: 0.5 的概率失去 1 000 元, 0.5 的概率得到 0 元;

概率分布 **d**: 确定性地失去 500 元。

此人的行为与期望效用理论不一致吗?

习题 3.9 证明:

$$(1) \sum_{t=0}^{\infty} \delta^t = \frac{1}{1-\delta}, (2) \sum_{t=0}^{\infty} \delta^{t-1} = \sum_{t=0}^{\infty} \delta^t, (3) \sum_{t=0}^T \delta^t = \frac{1-\delta^{T+1}}{1-\delta},$$

$$(4) \sum_{t=T}^{\infty} \delta^t = \frac{\delta^T}{1-\delta}, (5) \sum_{t=T}^S \delta^t = \frac{\delta^T - \delta^{S+1}}{1-\delta}.$$

习题 3.10 求 $\sum_{t=0}^{\infty} \delta^{2t}$ 。

习题 3.11 如果贴现因子是 $\delta = \frac{2}{3}$, 收益数列 (1, 2, 1, 2, 1, 2, 1, ...) 的值是多少?

习题 3.12 无限次地抛有三个面的骰子。如果第一个面落地, 你得到 1 元钱; 如果第二个面落地, 你得到 0 元钱; 如果第三个面落地, 你得到 2 元钱。假设骰子以第 x 面落地的概率是 π_x , 并且每次掷骰子的结果都是独立事件。假设这个骰子是均匀的, 贴现因子是 δ 。这一概率分布序列的期望值是多少? 如果财富的效用函数是 $u(x) = \ln x$, 这一概率分布的确定性等价是多少?

习题 3.13 在前一道习题中, 假设骰子是特制的。掷骰子的结果遵循马尔科夫过程, 对 $i, j \in \{1, 2, 3\}$, 转移概率为 π_{ij} 。设贴现因子是 $\frac{3}{4}$ 。 $\pi_{ii} = \frac{2}{3}$; 如果 $i \neq j$, 则 $\pi_{ij} = \frac{1}{6}$ 。回答前一题目中的问题。

习题 3.14 $\{x_t\}_{t=1}^{\infty}$ 为收益数列。要使 $\sum_{t=0}^{\infty} x_t \delta^{t-1}$ 取有限值, 这一数列必须满足什么条件?

习题 3.15 参数 α 和 γ 取什么值时, $\sum_{t=1}^{\infty} h(t) = \sum_{t=1}^{\infty} (1 + \alpha t)^{-\frac{\gamma}{\alpha}}$ 取有限值?

习题 3.16 如果所使用的是双曲线贴现, $h(t) = (1 + \alpha t)^{-\frac{\gamma}{\alpha}}$, 则收益数列 $(1, 2, 1, 2, 1, 2, 1, \dots)$ 的值是多少?

▶ 4

社会选择理论

4.1 公开招聘

这一部分所讨论的情况,是本书许多读者在职业生涯中很快会遇到的(或已经遇到的)情况:公开招聘师资。我们根据现实虚构一个政治科学系,它的师资均匀分布在五个专业中:美国政治(A)、比较政治学(C)、国际关系(I)、政治理论(T)和数理政治学理论与方法(F)。它所在的大学在财务上捉襟见肘,所以,院长只允许政治科学系增加一位教师。院长不愿意得罪系中任何一派,所以没有明确应该为哪个专业方向增加师资,只是说:“你们既然研究政治学,应该能相当好地解决这个问题。”在自己系里有过类似经历的读者,自然能领悟院长的挖苦之意。

同一专业领域中的教师,对于引进哪方面的人才,自然有相同的偏好。实际上,每个专业中的教师对这五个专业,都有自己的完全的和传递的偏好序。表 4.1 列出了他们各自的偏好序。

表 4.1

A	C	I	T	F
A	C	I	T	F
F	T	C	I	A
C	I	T	C	I
T	A	F	F	C
I	F	A	A	T

从表中可以看出,研究美国政治的人认为,最好是招聘研究美国政治的学者,最糟糕的莫过于招聘研究国际关系的学者;如此等等。系

主任得决定系里的招聘方式。他想到的第一个主意是,实行得票最多者胜出的规则(plurality rule),由全系教师进行投票。每个教师对自己最偏好的专业投票,得票最多者胜出。但是,系主任很快想到,这样做的结果是各个专业得票相等,打成平局。他决定实行两两投票多数通过(pairwise majority)的投票规则。每个专业都得与其他每个专业两两配对,投票决胜负。如果某个专业在所有这种两两投票中胜出,系里就为这个专业招聘师资。系主任自觉这是种公平的决策方式,就在全系大会上推行这一投票程序。各个专业首先对专业 A 和专业 C 投票。专业 C 得到专业 T 和专业 I 的支持。接着对专业 C 和专业 I 投票。I 以 3 比 2 的票数胜出。接着,专业 T 打败了专业 I;又先后打败了专业 F 和专业 A;最后却输给了专业 C。就是说,每个专业都在至少一轮的两两对决中落败。系主任采用的投票程序没有

产生任何结果。但是系主任看到,在所有的两两对决中,专业 A 不曾胜出过一次;专业 F 除了战胜过专业 A 外,输给了其他所有专业。现在,系主任至少知道,所增加的师资名额不应该给美国政治学和数理政治理论两个专业。

在无功而返后,系主任觉得,采用类似给大学橄榄球队排座次的打分系统也许管用。尽管有前面的挫折,他还是要求每个教师对所有专业进行排位。排位第一的专业得到 5 分,排位第二的专业得到 4 分,如此等等[这种程序被称为博尔达计分(Borda Count)]。系主任知道,如果每个人都按照自己的偏好投票,最终的排位依次是 C(17 分)、I(16 分)、T(15 分)、F(14 分)和 A(13 分)。

在拿到排位结果时,系主任惊讶不已。数理政治学家们看到了耍手段的机会,在投票时把专业 I 排在了第一位,把专业 C 排到了第五位。结果,专业 I 得到 18 分,专业 C 得到 16 分——这正是专业 F 所偏好的排位顺序。专业 C 被这种勾当激怒,要求重新投票。他们想在投票时,把专业 I 排在第五位,从而使专业 C 有机会胜出。但是,系主任意识到,如果重新投票,专业 I 会把专业 C 放到最后一位,借此报复;如果专业 T 在投票时也要花招,最终的胜出者甚至可能是专业 T。他宣布休会,给院长打了电话,把这个师资名额转给了经济学系。

这种集体选择问题落得如此下场并不使人意外。社会选择理论的一个基本结论是,集体选择过程或者限制了选择集,或者限制了偏好组合集,或者违反其他某个合理规范。在后面我们将看到,所有合理的集体决策机制都会遭遇到策略性的操纵,如在博尔达计分例子中专业 F 所做的勾当。

4.2 偏好加总规则

本节为正式分析偏好加总规则提供基本概念和思路。我们仅讨论有限个决策主体的情况,即 $N = \{1, 2, \dots, n\} (n > 2)$, 他们必须从集合 X 中选择一个结果。^① 我们的目的是介绍如何把每个人的偏好加总为集体偏好。第 i 个人在 X 上的偏好序为 R_i 。和第 2 章一样,个人的偏好序是完全的、反身的和传递的。我们用 $\mathcal{R}$ 表示由所有可能的完全的、反身的和传递的偏好序构成的集合。我们把所有 n 个人的偏好序,表示为 $\rho = \{R_1, R_2, \dots, R_n\}$, 称其为偏好组合。于是,由所有的这种组合构成的集合,就是 $\mathcal{R}^n$ 。集合 $\mathcal{B}$ 是由 X 上的完全的社会偏好序构成的集合。 $\mathcal{B}$ 中的偏好序未必具有传递性。

定义 4.1 偏好加总规则是函数 $f: \mathcal{R}^n \rightarrow \mathcal{B}$ 。

偏好加总规则只是个程序,它以个人偏好序集为基础,产生社会偏好序。我们用 R_i 表示个人偏好序,而 R 表示社会序。我们看本章开头的两两投票多数通过的投票规则。如果偏好序为 $xR_i y$ 的人数至少和偏好序为 $yR_i x$ 的人数相同,相应的偏好加总规则就为 xRy 。给定任意的偏好集,通过两两投票多数通过的投票规则,

^① 本章介绍完全信息模型,所以,“行动”、“政策”或“结果”等在这里含义相同。

总能得到完全偏好序,所以,这种投票程序满足上述定义。重要的是,这一定义并不要求偏好加总规则产生的结果具有传递性。实际上,在我们的政治学系中,两两投票多数通过的规则,产生的上一个循环序“ $TPIPCPT$ ”。

我们希望偏好加总规则满足哪些特征呢?也许最重要的特征是,它能产生最好的结果从而使这个集体能够真正地做出选择。换句话说,对所有的偏好组合,社会最大集 $M(R, X)$ 应该是非空集。第2章告诉我们,在偏好 R 满足非循环性或者传递性时,我们就能做到这一点。

定义 4.2 偏好加总规则 f 是传递的,如果对每个 $p \in \mathcal{R}^n$, f 所生成的社会序是传递的。

其次,偏好加总规则至少应具有最低程度的民主,而不是由某个人或者某个独裁者的偏好唯一地决定政策上的社会排序。

定义 4.3 偏好加总规则 f 是非独裁的,如果不存在某个 $i \in N$ 使得对每个 $p \in \mathcal{R}^n$ 和每个 $x, y \in X$, $xP_i y$ 意味着 xPy 。

这一条件相当严格。如果它被违反,“独裁者”就能在每种可能的偏好组合下,贯彻自己的偏好序。

偏好加总规则不应产生所有人都不接受的社会排序。在 x 和 y 之间,如果所有人都偏好 x ,则社会偏好应该反映这一排序。这一标准用意大利经济学家 Vilfredo Pareto 的名字命名,称为帕累托效率或帕累托最优。

定义 4.4 偏好加总规则 f 是弱帕累托最优的,如果对每个 $i \in N$, 对所有的 $x, y \in X$, $xP_i y$ 意味着 xPy 。

最后,对任意两种结果的社会偏好序,应该只决定于各个人对这两种结果的个人偏好序。在前述例子中,数理政治学家之所以能够操控系主任的投票程序的结果,一个原因是 C 和 T 之间的社会排序,决定于 F 对 F 、 A 和 I 的相对偏好。如果仅在 C 和 T 之间做选择,其他的偏好应该不产生任何影响。这一性质被称为无关选项的独立性(IIA)。

定义 4.5 偏好加总规则 f 独立于无关选项,如果对任意一对政策 $x, y \in X$ 和任意两个偏好组合 $p, p' \in \mathcal{R}^n$, 只要对每个 $i \in N$, $xR_i y$ 当且仅当 $xR'_i y$, 就有 xRy 当且仅当 $xR'y$ 。

所有这些特征都合情合理,并且其中每一个特征都很容易找出支持其成立的规范理由或者实际理由(IIA 的理由要弱一些)。但是,社会科学中最重要的结论之一是,加总规则不可能同时满足所有这些性质。Arrow 定理(1951)指出,产生传递性偏好、具有帕累托效率且满足 IIA 的唯一的加总函数是独裁。换句话说,只有当社会偏好是某个人(独裁者)的个人偏好时,社会偏好才具有个人偏好特征。

我们现在给出并证明 Arrow 定理。

定理 4.1 如果 X 是有限集且包含至少三个选项,则不存在具有传递性、非独裁性、弱帕累托最优性和独立于无关选项的偏好加总规则 $f: \mathcal{R}^n \rightarrow \mathcal{B}$ 。

这一结论的一个重要的但常常没有被正确理解的特点是,我们所寻找的是对 $\mathcal{R}^n$ 中所有可能的组合都有定义的加总规则。这一结论没有说,对每个偏好组合 $\rho \in \mathcal{R}^n$,满足IIA、弱帕累托最优性和非独裁性的加总规则都产生无传递性的偏好序;而是说,存在某个偏好组合,在这个偏好组合上,传递性被违反。一些人在表述Arrow定理时,常增加一个条件,称为“不受限制的定義域”,就是说, $\mathcal{R}^n$ 中每个偏好组合都是可能的。我们没有这么做,但是我们明确要求偏好加总规则的定義域是 $\mathcal{R}^n$ 。在证明Arrow定理前,我们还需要另一个定义。

定义 4.6 给定加总规则 f ,我们说集合 $W \subset N$ 对于 x 胜出 y 是半决定性的,如果在每个 $\rho \in \mathcal{R}^n$ 中,对所有的 $i \in W$,有 $xP_i y$,对所有的 $j \in W^c = N \setminus W$,有 $yP_j x$,且规则 f 产生 xPy 。如果在每个 $\rho \in \mathcal{R}^n$ 中,对所有的 $i \in W$,有 $xP_i y$,且规则 f 产生 xPy ,我们就说对于 x 胜出 y , W 是决定性的。如果对所有的 $x, y \in X$,集合 W 对于 x 胜出 y 是决定性的,我们就说 W 是决定性的。

我们首先证明在偏好加总规则满足某些条件时,决定性集合所具有的一个性质。在此基础上,我们给出Arrow定理的一个简便证明。

引理 4.1 如果 f 是个满足传递性的偏好加总规则,并且独立于无关选项且具有弱帕累托效率,如果对某个 $x, y \in X$, $W \subseteq N$ 对于 x 胜出 y 是半决定性的,则 $W \subset N$ 是决定性的。

证明:假设对于 x 胜出 y , $W \subset N$ 是半决定性的,并且在偏好组合 $\rho \in \mathcal{R}^n$ 中,对所有的 $i \in W$, $xP_i z$ 。假设还有另一个偏好组合 $\rho' \in \mathcal{R}^n$,它有以下特征:(1)对所有的 $i \in W$, $xP'_i yP'_i z$;(2)对所有的 $z \notin \{x, y\}$ 和所有的 $j \in W^c$, $yP'_j x$ 且 $yP'_j z$;(3)对所有的 $z \notin \{x, y\}$ 和所有的 $j \in W^c$, $xR_j z$ 当且仅当 $xR'_j z$ 。就是说,在个人对 x 和 z 的排序上, ρ 和 ρ' 相同。由于对于 x 胜出 y , W 是半决定性的,所以 $xP'y$ 。 f 具有弱帕累托效率,所以 $yP'z$; f 具有传递性,所以 $xP'z$ 。但是,由于(1) ρ' 没有明确 W^c 在 x 和 z 上的偏好,(2) ρ 和 ρ' 在 x 和 z 的排序上相同(就是说, $xR_j z$ 当且仅当 $xR'_j z$),所以, f 具有无关选项独立性的假设意味着 xPz 。因此,对于 x 胜出 z , W 是决定性的。这自然意味着对于 x 胜出 z , W 是半决定性的。同样的道理,对于 x 胜出 y , W 是决定性的。我们现在证明,对于 y 胜出 z , W 是决定性的。我们分析两个偏好组合 $\rho^0 \in \mathcal{R}^n$ 和 $\rho^+ \in \mathcal{R}^n$:在 ρ^0 中,对所有的 $i \in W$, $yP^0_i z$ 。在 ρ^+ 中,对所有的 $i \in W$, $yP^+_i xP^+_i z$;并且对所有的 $j \in W^c$, $zP^+_j x$ 且 $yP^+_j x$ 。假设对所有的 $j \in W^c$, $yR^0_j z$ 当且仅当 $yR^+_j z$ 。由于 W 对于 x 胜出 z 是决定性的,所以, xP^+z 。由于 f 是弱帕累托最优的,所以, yP^+x 。由于 f 具有传递性,所以, yP^+z 。 ρ^+ 只明确了 W 中的成员在 $\{y, z\}$ 上的偏好,且 ρ^0 和 ρ^+ 在 y 和 z 的排序上相同。因此,IIA意味着 yP^0z ,所以,对于 y 胜出 z , W 是决定性的。这意味着 W 对于 y 胜出 z 是半决定性的。利用这一事实并重复前面的步骤,我们就能看出, W 对于 y 胜出 z 是决定性的。综合所有这些结论,我们看到, W 是决定性的。□

就是说,如果对于某个两两比较,某个团体是半决定性的,则它就是决定性的。如果偏好加总规则满足IIA和弱帕累托标准,在某个两两比较中起决定作用的团体,

就在所有的两两比较中起决定性作用。我们现在证明 Arrow 定理——证明或者存在某个人,他是决定性的,或者整个集体不是决定性的。前一结论违背了非独裁性条件,后一结论违背了弱帕累托条件。这就是说,Arrow 的诸条件在逻辑上不兼容。

Arrow 定理的证明

设 X 是有限集,至少包含三个选项。为得到矛盾的结论,我们假设偏好加总规则满足传递性、非独裁性、弱帕累托效率和无关选项独立性。根据引理 4.1,对任意集合 $W \subset N$, 或者 W 是决定性的,或者不存在任何一对 $x, y \in X$ 使得 W 对于 x 胜出 y 是半决定性的。我们看两个不交集 $A, B \subset N$ 。设对任何 x 和 y , 它们都不是半决定性的(因此不是决定性的)。设 $C = N \setminus (A \cup B)$ 。由于 $n > 2$ 且任何单元素集合 $\{i\}$ 都不是决定性的,所以存在这样的三个集合 A, B 和 C 。我们分析偏好组合 $\rho \in \mathcal{R}^n$ 。设在偏好组合中,对所有的 $i \in A, xP_i yP_i z$; 对所有的 $j \in B, zP_j xP_j y$; 对所有的 $t \in C, yP_t zP_t x$ 。由于对任意一对选项, A 和 B 都不是半决定性的,所以, $zR^- x$ 且 $yR^- z$ 。由于 f 具有传递性,所以, $yR^- x$ 。这意味着集合 $A \cup B$ 对于 x 胜出 y 不是半决定性的,并且集合 $A \cup B$ 不是决定性的。这就是说,两个不是决定性的不交集的并,不是决定性的。由于 f 是非独裁的,所以没有任何单元素集合是决定性的。这就是说,没有任何团体是决定性的。这意味着 N 不是决定性的。这与 f 具有弱帕累托效率的前提相矛盾。□

本章开头的例子佐证了 Arrow 定理:两两投票多数通过的规则不具有传递性,博尔达计分规则不满足 IIA。再看一致通过的投票规则。对 x 和 y , 如果 xPy 当且仅当对所有的 $i, xR_i y$, 且对某个 $i, xP_i y$, 我们就说加总规则为一致通过规则。这一规则显然具有弱帕累托效率并满足 IIA, 因为这一规则根据各个人在 x 和 y 上的偏好, 在 x 和 y 之间做出选择。但它不具有传递性。为看到这一点, 我们分析表 4.2 中的个人偏好序。

表 4.2

1	2	3
y	z	z
z	y	x
x	x	y

很明显,一致通过规则意味着 xRy 且 yRz 。但是,这一规则也意味着 zPx 。

偏好加总规则的定义域是由所有可能的偏好组合构成的集合,所以,对某一偏好组合,可能依然存在着令人满意的加总偏好的方式。Arrow 定理并没有否定这一点。于是,对 Arrow 定理的一种反应是对偏好组合集施加限制,探讨是否存在在这一受限制的集合上满足规范公理的偏好加总规则。

一种常见的限制是单峰性。直观地说,单峰性要求结果按一定的顺序排列,从而使每个人的偏好序逐步上升至最偏好的结果,在此结果之后逐步下降。图 4.1 给出了政治科学系的偏好序。在图中,各专业领域按照“ACITF”的顺序排列。只有专业 I 和专业 T 有单峰偏好,他们的偏好序逐步上升至各自的理想结果,然后逐步下降。其他专业则有多峰偏好。例如,专业 A 在 A 和 F 上到达巅峰。喜欢钻研的读者会发现,我们无法通过调整结果的先后顺序而使所有偏好都有单峰。这就是说,我们的政治科学系的偏好组合不是单峰的。

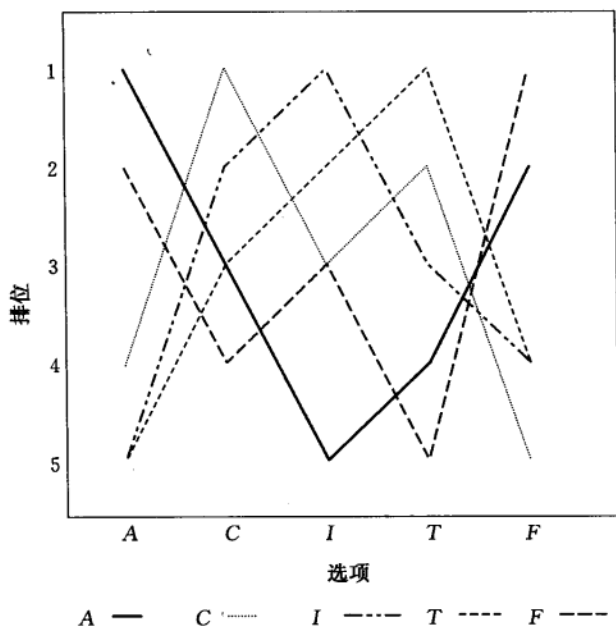


图 4.1 专业领域上的偏好

表 4.3

A	C	I	T	F
A	C	I	T	F
F	T	C	C	A
I	I	A	I	I
C	A	F	A	C
T	F	T	F	T

我们再看表 4.3 中的偏好组合。现在,如果把结果按照“TCIAF”(或“FAICT”)的顺序排列,则所有专业领域的偏好都有单峰,位于巅峰的结果正是他们自己的专业领域,见图 4.2。为给本节的下一个主要结论打下基础,我们分析两两投票多数通过规则生成的结果。现在,在所有的候选项中,专业 I 胜出。而且,两两投票多数通过的规则产生了

具有传递性的严格偏好序“IACFT”,正好与专业 I 的偏好相同。那么,多数通过的投票规则在我们的这个单峰偏好组合上产生了如此奇妙的结果,是否纯属偶然呢?答案是否定的。单峰性是多数通过规则具有传递性的充分条件(但不是必要条件)。

要表述和证明这一结论,我们需要引入一些符号。设 q 是个排序函数,定义域是结果集,它给每个结果赋予唯一的一个位置。正式地说, $q: X \rightarrow \{1, 2, \dots, |X|\}$ 是个单射和满射函数(或者说,是双射)。^①我们现在定义单峰性。

定义 4.7 给定集合 N 和有限选择空间 X , 偏好组合 $p \in \mathcal{R}^n$ 是单峰的, 如果存在某个双射 $q: X \rightarrow \{1, 2, \dots, |X|\}$ 使得对每个 $i \in N$, 存在某个 $t_i \in X$ 使得如

① 函数 $q: X \rightarrow X$ 是单射, 如果对每个 $y \in X$, 集合 $q^{-1}(y) = \{x \in X: q(x) = y\}$ 是单元素集合。它是满射, 如果对每个 $y \in X$, 存在某个 $x \in X$ 使得 $q(x) = y$ 。

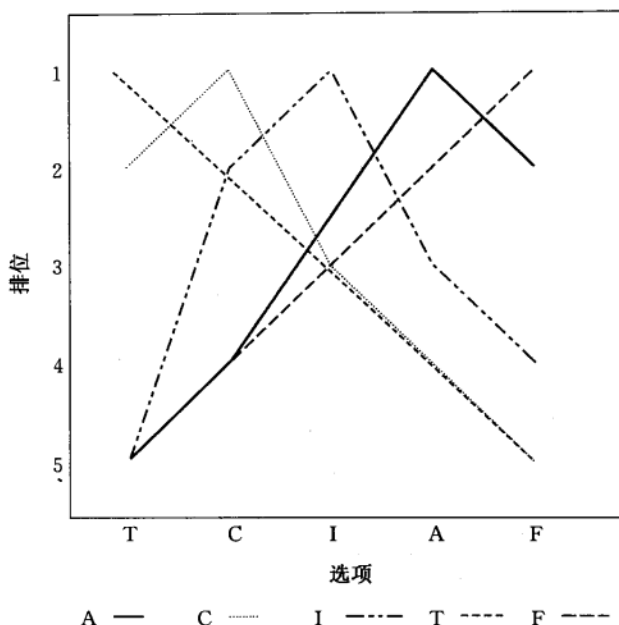


图 4.2 专业领域上的单峰偏好

055

果 $q(y) < q(t_i)$ 则 $t_i P_i y$ (如果 $q(x) < q(y) < q(t_i)$, 则 $t_i P_i y P_i x$), 如果 $q(t_i) < q(b)$ 则 $t_i P_i b$ (如果 $q(t_i) < q(b) < q(c)$, 则 $t_i P_i b P_i c$)。由所有的单峰偏好组合构成的集合被表示为 $S \subset \mathcal{R}^n$ 。

在这一定义中, 政策 t_i 是 i 的理想政策, 因为它是 $M(R_i, X)$ 中的唯一元素。随着 $q(x)$ 向上和向下偏离 $q(t_i)$, i 的偏好序下降。

现在给出正式的定理。

定理 4.2 给定 $\rho \in S$, 多数通过规则是传递的、弱帕累托最优的、IIA 和非独裁的。

它的证明很直接。我们先证明传递性。设 x, y 和 z 是三个选项。不失一般性, (利用定义 4.7) 设单峰序为 $q(z) > q(y) > q(x)$ 。给定 ρ , 设 $P(x, y; \rho)$ 是在 x 和 y 之间偏好 x 的所有人构成的集合。

单峰性对 x, y 和 z 上的偏好序施加了限制。例如, 设某个人在 x 和 y 之间偏好 x 。一旦施加了单峰性, 此人只能有一种偏好序, 即 $x P_i y P_i z$, 其他任何偏好序在 x 和 z 上有巅峰。所以, 我们知道, 在 x 和 z 之间, 此人偏好 x 。因此, 在 x 和 y 之间偏好 x 的所有人, 在 x 和 z 之间也偏好 x , 于是, $P(x, y; \rho) \subseteq P(x, z; \rho)$ 。这里请你证明单峰性产生下面这些条件:

$$P(x, y; \rho) \subseteq P(x, z; \rho)$$

$$P(z, x; \rho) \subseteq P(y, x; \rho)$$

$$P(z, y; \rho) \subseteq P(z, x; \rho)$$

$$P(x, z; \rho) \subseteq P(y, z; \rho)$$

多数通过规则的一个特征是,给定某一群体,如果其任意子集在一对选项上是决定性的,则这一群体在这对选项上是决定性的。如果一个群体构成多数,它所在的所有群体由于都比它大,因此也构成多数。因此,单峰性所施加的限制意味着:

$$\text{如果 } xPy, \text{ 则 } xPz \quad (4.1)$$

$$\text{如果 } zPx, \text{ 则 } yPx \quad (4.2)$$

$$\text{如果 } zPy, \text{ 则 } zPx \quad (4.3)$$

$$\text{如果 } xPz, \text{ 则 } yPz \quad (4.4)$$

要证明传递性,我们需要证明以下六个陈述彼此之间不矛盾。其中每个陈述或者直接得自上面某个条件,或者它的前提不成立。^①

- (1) 如果 xPy 且 yPz , 则 xPz 。这一陈述直接得自(4.1)式。
- (2) 如果 xPz 且 zPy , 则 xPy 。这一陈述的前提与(4.3)式矛盾。
- (3) 如果 yPx 且 xPz , 则 yPz 。这一陈述得自(4.4)式。
- (4) 如果 yPz 且 zPx , 则 yPx 。这一陈述得自(4.2)式。
- (5) 如果 zPx 且 xPy , 则 zPy 。这一陈述的前提与(4.2)式矛盾。
- (6) 如果 zPy 且 yPx , 则 zPx 。这一陈述得自(4.3)式。

我们穷尽了所有可能性,证明了在选择空间中有三个选项时,如果偏好是单峰的,则多数通过规则具有传递性。在此基础上的扩展并不很难。这里请你证明多数通过规则是弱帕累托最优的、IIA 和非独裁的。

4.3 集体选择

我们从加总偏好序应具有的性质开始了我们的讨论,但我们最终想得到的是在给定的偏好加总规则下,由最优政策构成的集合。类似个人选择,我们假定社会的选择来自于由加总偏好序决定的最优集。在社会选择背景中,我们称最优选择集为核。

定义 4.8 给定 $X, \rho \in \mathcal{R}^n$ 和偏好加总规则 f , 核为 $C_{f(\rho)}(X) = M(f(\rho), X)$ 。

如果 X 是有限集且集体偏好是完全的和传递的,应用定理 2.1,我们能证明,核是非空集,社会选择能够做出。但是,Arrow 定理指出,偏好加总规则的传递性并不总能得到满足。这一性质一旦得不到满足,核就可能是空集。

在多数通过规则下,如果核存在,它的结果就被称为 Condorcet 赢家——用 Marquis de Condorcet 的名字命名,他是最早正式研究投票程序的性质的人之一。

① 一个陈述的前提(即“如果”部分)如果不成立,则这一陈述成立。

但是,在单峰偏好下多数通过规则的核非空的结论,一般认为是 Duncan Black 提出的。

定理 4.3 设 $n > 2$ (为奇数), $p \in S$ 且 $f(\cdot)$ 为两两投票多数通过规则,则:

$$C_{f(p)}(X) = \{t_m : |j \in N \setminus m: q(t_j) \leq q(t_m)| = |k \in N \setminus m: q(t_k) \geq q(t_m)|\}$$

就是说,核是中间投票人的理想点。

这个定理的证明类似单峰偏好下多数通过规则的传递性的证明,但更简单。设 t_m 是中间投票人的理想点,于是, $\frac{n}{2}$ 个人的理想点 t_i 满足 $q(t_i) \leq q(t_m)$, 另外 $\frac{n}{2}$ 个人的理想点 t_i 满足 $q(t_i) \geq q(t_m)$ 。我们定义 $P(t_m, x)$ 为在 t_m 和 x 之间偏好 t_m 的人构成的集合。如果 $q(x) < q(t_m)$, 单峰性意味着 $P(t_m, x)$ 肯定包括中间投票人和理想点 t_i 满足 $q(t_i) > q(t_m)$ 的所有人。这一集合构成多数,因为它包含至少 $\frac{n+1}{2}$ 个人。如果 $q(x) > q(t_m)$, $P(t_m, x)$ 就包括中间投票人和理想点 t_i 满足 $q(t_i) < q(t_m)$ 的所有人。因此,对所有的 $x \neq t_m$, $t_m P x$, 所以, t_m 是最优集的唯一元素。

这一结论指出,如果偏好是单峰的,则多数通过规则的核存在。在这种情况下,“多数人决定”可能是做出集体选择的合理的方式。

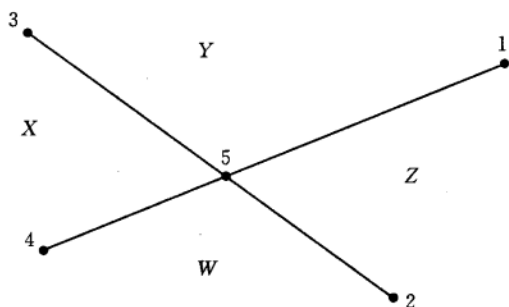


图 4.3 在两维空间中的 Condorcet 赢家

但是,局限于单峰偏好可能未必始终合理。例如,政策集可能为多维空间。此时,我们很难把单峰性推广至这种情况中。图 4.3 给出了两维空间中五个投票人的理想点。每个人都有圆形无差异曲线,因此,在任何两两比较中,他都偏好距离自己的理想点最近的选项。这些无差异曲线对应欧几里得偏好。点 5 是多数规则中的核点或者说是 Condorcet 赢家,因为在这一政策空间中,多数人偏好这个点。为证明这一结论,我们证明对其他任何政策,至少存在三个反对者。假设发生了由这个政策向区域 W 中任意政策的转变。很明显,投票人 5 反对这种变化;投票人 1 和 3 亦如此。就是说,相对于区域 W 中任何的政策,5 的理想点是多数人偏好的政策。投票人 1、2 和 5 反对区域 X 中的政策;投票人 2、4 和 5 投票反对区域 Y 中的

政策;投票人 3、4 和 5 投票反对区域 Z 中的政策。

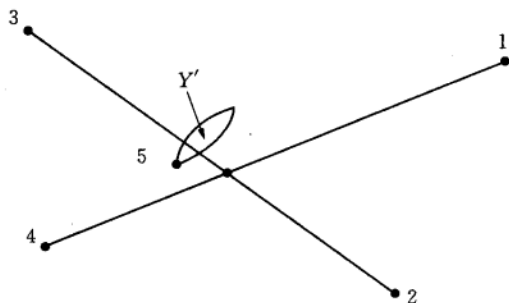


图 4.4 无 Condorcet 赢家

投票人 5 的理想点是核点,因为在任意两个选项上,投票人 5 都是中间投票人:在 x 和 y 之间,如果他偏好 x ,则至少存在另外两个投票人亦如此。由于在比较其他人的理想点时,这一结论亦适用,所以,投票人 5 的理想点肯定位于连接相对立的两个理想点的直线上(2—3 和 1—4)。这一条件与 Plott(1967)条件有密切关系。下一节将讨论这一条件;在这里,我们要看到,这一条件非常脆弱。图 4.4 指出,对这些直线的交点的任何偏离,都会破坏多数通过规则的核。首先,这个交点不可能是核点,因为投票人 3、4 和 5 更偏好投票人 5 的理想点。其次,投票人 5 的理想点不可能是 Condorcet 赢家,因为存在点集 Y' ,投票人 1、2 和 3 更偏好这个集合中的点。

鉴于多数通过规则的核的存在性条件(更准确地说,非空性条件)如此苛刻,我们不禁要问,多数通过规则能否通过剔除某些不合理的结果,缩小结果集。在本章开头的例子中,系主任认为不应该在美国政治和数理理论方向增加师资——因为这两个专业屈居另三个专业之后。因此,多数通过规则剔除了这两个选项,但在前三个选项上产生了循环偏好。就是说,C、I 和 T 构成了顶端循环集(top cycle set);这个集合中的选项优于集合外的所有选项,但是,在这一集合内,多数通过的加总规则不具有传递性。多数通过规则虽未必能生成核,但它也许能生成一个很小的顶端循环集。遗憾的是,这一点亦未必能做到。Richard McKelvey(1976)指出,在诚实投票(面对任何一对选项,每个人投票给自己偏好的选项)和欧几里得偏好下,顶端循环集或者是核或者是整个选择集。下一节将正式探讨 McKelvey 的结论,这里只用图 4.5 说明其中的道理。在图中,三个投票人都有二次方偏好,有一个初始现状政策 a 。要理解 McKelvey 的结论,我们只需证明,两两投票多数通过规则能够导致从 a 点到政策空间中任意一点的变化。首先,和 a 点相比, b 点能得到多数支持,因为投票人 1 和投票人 2 偏好 b 。在 c 和 b 之间,投票人 2 和投票人 3 偏好 c 。在 d 和 c 之间,投票人 1 和投票人 3 偏好 d 。这一过程不断持续下去,于是,相对于当前的现状,多数人偏好的政策集变得越来越大。我们就能越来越远地偏离各个投票人的理想点。最终,在 e 和 d 之间,投票人 1 和投票人 2 偏好更远

的点 e 。从 e 点开始,我们可能回到 a 点(在 a 和 e 之间,投票人一致偏好 a),或者到达更偏远的点。

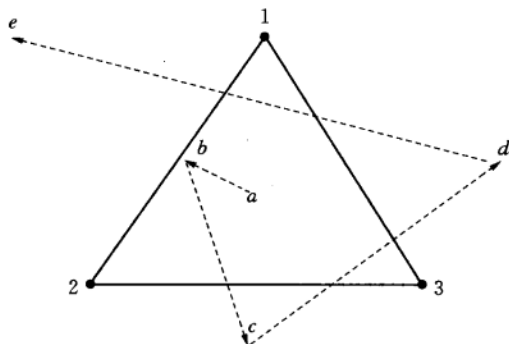


图 4.5 McKelvey 定理

这一偏好加总理论算不上是一种实证分析方法,因为它没有提供清晰的经验实证结论。一些人从 McKelvey 的模型中引申出政治混沌的结论:在有循环偏好时,选择是不稳定的。这种理解过于简单;它给一个没有任何结论的模型赋予了一个实证结论。更合理的解释是,McKelvey 模型说明了分析政治制度的影响的重要性——我们毕竟是在政治制度内做出集体选择的。所以说,一个仅以偏好和偏好加总规则为出发点的模型,常常显得有些无力。后面介绍的非合作博弈论,能使我们构造出在经验实证意义上更重要的集体选择模型。

对 Plott 条件和 McKelvey 定理的正式分析*

在本节中,我们在多数通过规则和多维度偏好背景中,分析 Plott 和 McKelvey 的结论。我们的分析建立在以下社会选择环境基础上。

条件 4.1 $X \subset \mathbb{R}^d$ 是凸的集,在 X 上每个人都有严格凸的连续偏好。

如果 X 是紧集,根据定理 2.3 和定理 2.4,每个人在 X 中都有唯一的理想点 y_i 。但是,这里并不假定 X 是紧集,而假定每个人都有理想点。严格凸偏好要求上半集(upper contour set)是严格凸集。如果将偏好限制为欧几里得偏好,即 $u_i(x) = -\|x - y_i\|$,我们就能精确描述政策变化所导致的效用变化。

定义 4.9 如果偏好为欧几里得偏好,则对任意的 $x \in X$,梯度向量 $\nabla u_i(x) = y_i - x$ 。

这样的梯度向量是个有方向的量或者说是个有方向的线段,指向第 i 个人希望政策在点 x 上的变化方向。对一般效用函数而言,点 x 上的梯度向量只是偏导数向量在点 x 的取值。

定义 4.10 给定效用函数 $u: \mathbb{R}^d \rightarrow \mathbb{R}$,梯度向量为 $\nabla u_i(x) = (\partial u(x)/\partial x_1, \partial u(x)/\partial x_2, \dots, \partial u(x)/\partial x_d)'$ 。

在 Plott 条件中,用到了“配对”(pairing)一词。给定有限集合 A , 映射 $p: A \rightarrow A$ 如果是双射,则为配对。就是说, A 中的每个点 i 与 A 中唯一一个点 j 配对。现在给出 Plott 条件。

定义 4.11 在有欧几里得偏好的空间模型中,我们说 Plott 条件在政策 $x \in X$ 上得到满足,如果在集合 $L = \{j \in N: y_j \neq x\}$ 上存在配对函数 $p(\cdot)$ 使得对每个 $i \in L$, 存在 $\lambda_i > 0$ 使得 $\nabla u_i(x) = -\lambda_i \nabla u_{p(i)}(x)$ 。

如果 Plott 条件在点 x 上得到满足,则理想点异于 x 点的每个人,能与另一个人配对,从而使得这一配对中的两个人希望政策向完全相反的方向变化。由于政策变化的每个支持者都对应着一个反对者,所以不可能出现多数联盟从而推翻政策 x 。下面的结论概括了在有欧几里得偏好的空间模型中,Plott 条件和多数通过规则的核之间的关系。

定理 4.4 在有欧几里得偏好的空间模型中,在人数 n 为奇数时,位于 X 的内点集中的 x 位于核 $C_{f(p)}(X)$ 中当且仅当在点 x , Plott 条件得到满足。^①

显然,在一般情况下,Plott 条件得不到满足。假设在某个 x 上,Plott 条件得到满足。则对所有的 i , $y_i - x = -\lambda_i(y_{p(i)} - x)$ 。如果我们扰动 $y_{p(i)}$, 使它位于另一向量上,则这一条件在 x 点不再成立。更准确地说,如果 $\mathbb{R}^n$ 是在选择空间 $\mathbb{R}^d$ 有欧几里得偏好的 n 个人的理想点构成的空间,则使得 Plott 条件在某个点 $x \in \mathbb{R}^d$ 上得到满足的 $\mathbb{R}^n$ 的子集极小。它不包含任何开子集,因此其内点集为空集。如果随机选择一个偏好组合,则选到在某个点上满足 Plott 条件的偏好组合的概率为零。

虽然有核点的偏好组合集很小,但是,有核点的每个偏好组合都任意地接近有核点的其他偏好组合。在习题中,我们请你证明,如果扰动某个有核点的偏好组合,则由此生成的新的偏好组合的核点(如果存在核点的话),任意地接近原来的核点。^②

核一般为空集。那么,是否存在政策空间的其他某个子集,拥有理想的规范特征并且能产生合理的结果呢? 顶端循环集是一种可能性。

定义 4.12 给定集合 X 、偏好组合 $p \in \mathcal{R}^n$ 和偏好加总规则 f , 顶端循环集为 $T_{f(p)} = \{x \in X: \forall y \in X \setminus x, \exists \{a_0, \dots, a_t\} \subset X \text{ 使得 } a_0 = x, \text{ 且对 } t < \infty, a_t = y, \text{ 且 } \forall z \leq t, a_{z-1} P a_z\}$ 。

顶端循环集是从其他任何点出发,通过有限的严格偏好链,能够到达的点构成的集合。就是说,如果 $x \in T_{f(p)}$, 则对每个 $y \in X \setminus x$, 我们能够找到有限个政策 $\{a_1, a_2, \dots, a_t\}$ 使得 $x P a_1 P a_2 P \dots P a_t P y$ 。下面的结论指出,或者 Plott 条件得到满足或者顶端循环集覆盖了整个政策空间(McKelvey 证明了这一结论)。

定理 4.5 在空间模型中,或者 $C_{f(p)}(X)$ 是非空集或者 $T_{f(p)} = X$ 。

① 点 x 位于 X 的内点集中,如果存在一个包含于 X 中的开球 $B(x, \epsilon)$ 。见数学附录。

② 用可微性代替欧几里得偏好假设,能得到更一般的结论。

上面两个定理意义重大。在有欧几里得偏好的空间模型中,除非刀锋状的 Plott 条件在某个政策上得到满足,否则任意政策都能够通过有限的严格偏好链,被其他任意一个政策达到。

4.4 操控选择函数

前一节中的讨论表明,多数通过规则常常无法为社会选择提供足够的指导。即使这一规则的核存在,人们也可能没有动机诚实地披露自己的偏好。Gibbard (1973)和 Satterthwaite(1975)指出,所有合理的社会选择函数,包括多数通过规则,都容易被投票人策略性地操控。在本章开始时,我们看到数理政治学家们通过虚假揭示自己的偏好,操控博尔达计数程序。在投票中,类似的操控可能发生。假设在某个投票表决中,首先就 x 与 y 投票,然后就胜出方案与 z 投票。假设在多数通过的规则下,有 xPy 、 zPx 和 yPz 。如果所有人都按照自己的真实偏好投票,则 x 打败 y ,却输给 z 。但是,在 x 与 y 之间偏好 x 且在 y 与 z 之间偏好 y 的投票人,有动机在第一轮投票中投票给 y (即虚假地揭示自己在 x 与 y 上的偏好),于是, y 可能在第一轮中胜出,并在第二轮投票中打败 z 。

在给出 Gibbard-Satterthwaite 定理前,我们先给出几个定义。

定义 4.13 社会决策函数是满射函数 $G: \mathcal{R}^n \rightarrow X$, 给定一个偏好组合,它产生一个结果。

给定一个偏好组合,社会决策函数产生一个结果。要求 G 是“满射”,是要求每个结果对应一个偏好组合。我们来定义操控。

定义 4.14 在偏好组合 ρ 上, $G(\rho)$ 是可操控的当且仅当存在某个 i , 存在另一偏好组合 $\rho' = \{R_1, \dots, R_{i-1}, R'_i, R_{i+1}, \dots, R_n\}$, 使得 $G(\rho') P_i G(\rho)$ 。如果在任何的 ρ 上, $G(\rho)$ 都是不可操控的,则说 G 是不可操控的。

给定一个社会决策函数,如果有人能通过假报自己的偏好,把结果改变为自己所偏好的结果,这一函数就是可操控的。在这一定义中,投票人 i 通过假报自己的偏好为 R'_i 而不是真实的 R_i , 把结果从 $G(\rho)$ 转变为他所偏好的 $G(\rho')$ 。我们给出 Gibbard-Satterthwaite 定理。

定理 4.6 如果存在三个以上的选项,并且 G 是不可操控的,则存在独裁者(就是说,存在某个 i , 对所有的 $x \in X \setminus G(\rho)$ 和所有的 $\rho \in \mathcal{R}^n$, $G(\rho) P_i x$)。

首先给出这一定理的证明思路。假设 G 是不可操控的。则应用 G 到选择集 X , 得到最偏好的结果;再应用 G 到 X 的剩余元素上,我们得到次佳结果。如此进行下去,我们就能构造出具有传递性和满足 IIA 的偏好序。由于这一偏好序满足弱帕累托标准,所以根据 Arrow 定理,存在独裁者。

现在给出详细证明。假设 G 是不可操控的。第一步,设偏好组合 ρ 使得存在集合 B , 对所有的 i , 对所有的 $x \in B$ 和所有的 $y \in X \setminus B$, $x P_i y$ 。就是说,和集合 B 以外的所有选项相比,所有人都偏好 B 中的选项。我们的第一个结论是,社会的决策

是这一“最好的”集合中的元素,即 $G(\rho) \in B$ 。假设不是如此。由于 G 是满射,所以,我们能找到另一偏好组合 ρ' 使得 $G(\rho') \in B$ 。我们于是能够构造出一个选项列:

$$y_0 = G(\rho)$$

$$y_1 = G(\rho | R'_1) = G(R'_1, R_2, \dots, R_n)$$

$$y_i = G(\rho | R'_1, \dots, R'_i) = G(R'_1, \dots, R'_i, R_{i+1}, \dots, R_n)$$

$$y_n = G(\rho')$$

设 k 是使得 $y_k \in B$ 的最小整数。由于在 B 内部的一切结果和 B 外部的一切结果之间,投票人 k 偏好前者,所以,通过报告偏好为 R'_k 而不是 R_k ,他得到了一个更好的结果。所以, G 并不是不可操控的。因此, $G(\rho) \in B$ 肯定成立。

第二步,给定 ρ ,构造加总规则或者说对所有的政策,进行社会排序。设排位最高的元素是 $x_1 = G(\rho)$ 。我们把 x_1 放到每个人的偏好序的最底部,借此产生 ρ_2 。排在次高位上的元素便是 $x_2 = G(\rho_2)$ 。根据步骤一, $x_2 \neq x_1$ 。持续进行这一过程,直至所有选项都得到排列。

不难看出,由此方式得到的偏好序,满足弱帕累托标准——在每个阶段上,决策规则为所构造的偏好组合从“最好的”集合中选择一个元素。现证明以此方式得到的加总规则满足 IIA。假设它不满足 IIA,则存在两个偏好组合 ρ 和 ρ' 和两个选项 $x, y \in X$, 使得 $xR_i y$ 当且仅当对所有的 i , $xR'_i y$, 并且 $xf(\rho)y$ 且 $yf(\rho')x$ 。设 $\rho(x, y)$ 是个偏好组合,它不同于 ρ 的唯一地方是 x 和 y 被移到了每个人的偏好序的顶端。我们来证明 $G(\rho(x, y)) = x$ 。假设这一结论不成立。我们就不断地把各个选项移到底部——借此形成新的偏好组合 $\hat{\rho}$ ——直至 $G(\hat{\rho}) = x$ 。我们看序列 $y_i = G(\rho(x, y) | \hat{R}_1, \dots, \hat{R}_i)$, y_i 是把前 i 个人换到这一新的偏好组合中时所产生的社会决策。在这里, $y_n = G(\hat{\rho}) = x$ 且 $y_0 = G(\rho(x, y)) = y$, 因为根据步骤一, $G(\rho(x, y)) \in \{x, y\}$ 。和前面一样,设 k 是使得 $y_k \neq y$ 的最小整数。如果 $y_k = x$ 且 $xP_k y$, 则通过报告 $\hat{R}_k$ 而不是 R_k , 就能够操控 G 。如果 $yP_k x$, 则通过报告 R_k 而不是 $\hat{R}_k$, 亦能操控 G 。如果 $y_k \neq x$, 则分析使 $y_j \in \{x, y\}$ 的最小的 $j > k$ 。利用和前面相同的逻辑,我们能够证明投票人 j 能够操控 G 。于是,我们就得到了与 G 不被操控的前提相矛盾的结论,并证明了 $G(\rho(x, y)) = x$ 。

现在看序列 $z_i = G(\rho(x, y) | R'_1(x, y), \dots, R'_i(x, y))$ 。偏好序不满足 IIA 的假设,意味着 $z_n = G(\rho'(x, y)) = y$ 且 $z_0 = G(\rho(x, y)) = x$ 。和前面一样,肯定存在某个人,在 x 和 y 之间他偏好 y , 并且为把结果从 x 变为 y , 他会假报自己的偏好为 $R'(x, y)$ 而不是真实的 $R(x, y)$; 或者存在某个人,在 x 和 y 之间他偏好 x , 并且他会假报自己的偏好为 $R(x, y)$ 而不是真实的 $R'(x, y)$ 。就是说, G 是可操控的。这一与前提矛盾的结论意味着前面得到的偏好序肯定满足 IIA。根据 Arrow 定理,在 f 上,肯定存在独裁者,因而在 G 上,肯定存在独裁者。

Gibbard-Satterthwaite 定理虽然在结论上相当悲观,却对研究政治行为有重要

意义。其最重要的意义也许是,在政治中策略性行为无所不在:“防范策略性行为”的机制常常并不存在。下一章开始研究政治行为的策略模型。

4.5 习题

习题 4.1 两个投票人在五个选项上,有如下偏好序

1	2
A	E
B	B
C	D
D	A
E	C

- (1) 每个选项得到的博尔达分数是多少?
- (2) 投票人 1 如何通过假报自己的偏好,得到更好的结果?
- (3) 投票人 2 如何通过假报自己的偏好,得到更好的结果?
- (4) 是否存在某个(不一定真实的)偏好陈述组合,两个投票人没有激励在事后改变自己的陈述?

习题 4.2 证明:给定 $\rho \in S$, 多数通过规则具有传递性、弱帕累托效率、IIA 和非独裁性。

习题 4.3 (Austen-Smith 和 Banks, 1999) 我们说偏好加总规则 f 满足投票人主权, 如果对所有的 $x, y \in X$ (其中 $x \neq y$), 存在 $\rho \in \mathcal{R}^n$ 使得 xPy 。我们说偏好加总规则 f 是单调的, 如果对所有的 $x, y \in X$, 对所有的 $\rho, \rho' \in \mathcal{R}^n$, 以下命题成立: 如果 (1) $xP_i y$ 意味着 $xP'_i y$, (2) $xR_i y$ 意味着 $xR'_i y$, 和 (3) xPy , 则肯定有 $xP' y$ 。证明: 如果 f 是弱单调的并且满足投票人主权, 则 f 具有弱帕累托效率。

习题 4.4 三个投票人要通过两两投票多数通过规则选择一项政策。如果有三个选项, 所有人的偏好都是严格的, 并且每个人有不同于其他两人的偏好序, 则在全部的可能的偏好组合中, 有多大比例的偏好组合会导致 Condorcet 赢家? (如果其中两人有相同的偏好序, 则它们最偏好的结果肯定是 Condorcet 赢家。为什么?)

习题 4.5 三个投票人都有两维空间上的欧几里得偏好, 其理想点分别为 $(-1, 0)$ 、 $(0, 1)$ 和 $(1, 0)$ 。

- (1) 构造一个规则以便从 $(0, 0)$ 到达 $(2, 2)$;
- (2) 构造一个规则以便从 $(0, 0)$ 到达 $(5, 5)$;
- (3) 构造一个规则以便从 $(0, 0)$ 到达 $(-5, -5)$ 。

习题 4.6 证明: 如果 $\rho \in \mathbb{R}^n$ 是个理想点组合且在某个点 $x \in \mathbb{R}^d$ 上 Plott 条件得到满足, 则对每个 $\varepsilon > 0$, 存在组合 $\rho^\varepsilon \in B(\rho, \varepsilon)$, 在任何点上, Plott 条件都得

不到满足。

习题 4.7 证明:如果 $\rho \in \mathbb{R}^n$ (n 为奇数) 是个理想点组合, 在某个点 $x \in \mathbb{R}^d$ 上 Plott 条件得到满足, 则对每个 $\varepsilon > 0$, 存在 $\delta > 0$ 使得如果 $\rho^\delta \in B(\rho, \delta)$ 并且在某个点上, Plott 条件在组合 ρ^δ 上得到满足, 则在组合 ρ^δ 上, Plott 条件在某个点 $x' \in B(x, \varepsilon)$ 上得到满足。

鄧
平
船
解
學
PDG

▶ 5

标准式博弈

节目开始 12 分半钟后,警官罗根和布里斯科就抓到了两名杀人嫌犯。地区检察官亚当·希夫指示助理检察官杰克·马考伊向两名疑犯分别开出下面的条件:

- 你如果认罪并提供同伙实施一级谋杀的证据而你的同伙不认罪的话,你将以非法携带武器的罪名被判入狱一年。如果你的同伙认罪,你们都以二级谋杀罪名各被判入狱 8 年。

- 你如果不认罪而你的同伙成为检方证人的话,你将以一级谋杀罪被判入狱 25 年甚至终身监禁。他如果也不认罪,你们将以故意杀人罪名各被判入狱 4 年。

我们假设每坐牢一年,嫌犯减少一单位效用。表 5.1 给出了在所有可能的结果中,每个嫌犯获得的收益。表中的行给出了嫌犯 1 的行动,列给出了嫌犯 2 的行动。每对数字表示在一种行动组合中,嫌犯 1 和 2 得到的收益。

两个嫌疑犯处在策略性环境中,因为嫌犯 1 的行动所导致的结果,取决于嫌犯 2 的选择;反之亦然。那么,这两个嫌犯应该怎么做呢?就集体利益而言,他们不应该认罪。因为,如果都不认罪的话,两人坐牢时间共为 8 年,远低于其他结果下的

表 5.1 囚徒两难

1\2	不认罪	认罪
不认罪	-4, -4	-25, -1
认罪	-1, -25	-8, -8

坐牢时间。但是,除非这两个嫌犯已达成有约束力的协议,否则他们各自对个人利益的追求,会使这一结果无法实现。假设嫌犯 1 不认罪,嫌犯 2 意识到,自己如果认罪,就能得到更好的结果:入狱时间从 4 年减少到 1 年。事实上,两个嫌犯都认识到,不管对方采取什么行动,认罪都能给自己带来好处。于是,两个人都选择认罪,坐牢时间共为 16 年。在这里,个人理性导致了社会意义上更差的结果(这里的社会指的是由这两个嫌犯构成的社会;而检察官和警察更喜欢这个结果)。^①

这便是著名的“囚徒两难”博弈。在这个博弈中,我们能直接推断出理性参与

① 在这一剧集中,接下来的情节是:由于技术细节上的纰漏,纽约西区法院推翻了嫌犯的认罪证词。在剧集的最后,当马考伊给自己倒了一杯 18 年陈苏格兰威士忌时,希夫做了一番精辟点评。(美国电视连续剧 *Law and Order* 中的一集。——译者注)

者的策略。但是,在其他的策略性环境中,做出这样的预测并非易事。我们来分析一个“抓捕恐怖分子”博弈。联邦调查局(FBI)和中央情报局(CIA)负责调查和抓捕恐怖分子。有两类恐怖分子,分别为核心分子和行动人员。两个机构更想抓到的是核心分子;但他们都不想一无所获。抓捕核心分子需要两个机构协作,都投入资源到联合抓捕行动中。而要抓捕恐怖组织的行动人员,只需两个机构独自展开调查。表 5.2 给出了两个机构在决定抓捕核心分子或行动人员时,所处的策略性环境。

表 5.2 抓捕恐怖分子

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	2, 2	0, 1
抓捕行动人员	1, 0	1, 1

在这一矩阵中,行表示 FBI 的策略(抓捕核心分子或抓捕行动人员),列表示 CIA 的策略。如果抓到核心分子,两个机构各得到 2 个单位的效用;一个机构如果抓到行动人员,得到 1 个单位的效用;如果一无所获,得到 0 个单位的效用。我们先看 FBI 的决策。和囚徒两难博弈不同,在这里,FBI 的最优选择决定于 CIA 的选择。如果 CIA 在抓捕核心分子,则提供合作能使 FBI 得到 2 个单位的效用,否则它只能独自抓捕行动人员而得到 1 个单位的效用。如果 CIA 在抓捕行动人员,则抓捕核心分子的行动只能使 FBI 得到 0 个单位的效用。所以说,FBI 做出怎样的选择,决定于它对 CIA 所做出的选择的推断。对 CIA 来说,亦是如此。那么,每个机构对另一机构的行动,形成怎样的推断才是合理的呢?在对策略性行为的研究中,John Nash 为理性均衡行为的分析提供了重要工具。在 Nash 的理论中,每个机构所选择的策略,都是对另一机构的策略的“最优反应”。如果两个机构都以此方式做出选择,最终的结果便是对最优反应的最优反应,没有哪个机构有动机改变自己的策略。由于这种策略组合产生稳定的行为预测,因此被称之为均衡——为纪念 Nash(1950b)的贡献,它被称为 Nash 均衡。

要确定某个策略组合是否是 Nash 均衡,只需检验在这一策略组合上,没有哪个机构能通过单方面偏离到另一选择上而得到更高的效用水平。两个机构都去抓捕核心分子的结果是 Nash 均衡吗?如果 CIA 抓捕核心分子,FBI 的最优反应是抓捕核心分子。同样地,如果 FBI 在抓捕核心分子,CIA 的最优反应也是抓捕核心分子。因此,两个机构都抓捕核心分子的行为构成 Nash 均衡。但是,这不是这一博弈的唯一 Nash 均衡。如果 CIA 决定抓捕行动人员,FBI 的最优选择也为抓捕行动人员。由于在 FBI 抓捕行动人员时,CIA 抓捕行动人员能给自己带来更好的结果,所以,两个机构都抓捕行动人员的行为构成 Nash 均衡。^①Nash 的解没能提

① 还存在第三个均衡。在这个均衡中,每个机构都以 0.5 的概率抓捕核心分子,以 0.5 的概率抓捕行动分子。在下文中,我们将探讨这种可能性。

供唯一的均衡结果,但是,它排除了一些行为。例如,一个机构抓捕核心分子而另一机构抓捕行动人员的结果就不是 Nash 均衡——抓捕核心分子的机构,如果转而抓捕行动人员,就能得到更高效用。而抓捕行动人员的机构也愿意改变自己的选择。

上述例子尽管简单,但是 Nash 的均衡概念却是个相当有力的分析工具,能够分析大量博弈中的行为。所以,在本章中(和本书中),我们将推广和发展 Nash 均衡概念。

5.1 标准式博弈

用博弈论分析政治现象的第一步是表述策略性环境。这里介绍最简单的表述形式:完全且完美信息标准式博弈。它包括以下构成因素:

(1) 参与者:设 N 为参与者集合。我们用符号 $i \in N$ 表示任意一个参与者,用符号 $-i \in N$ (读为“非 i ”)表示 i 以外的所有参与者。

(2) 纯策略:某个参与者的一个纯策略是他的一个行动方案,如前面例子中的“认罪”或“抓捕行动人员”。在一次性博弈中,如前面的例子,一个策略就是一个行动。而在多阶段博弈中,一个策略是一个函数,它根据此前所有阶段发生的事件,规定每个阶段中要采取的行动。在标准式博弈中,每个参与者有一个纯策略集合:参与者 $i \in N$ 的策略集为 S_i 。 $s_i \in S_i$ 为参与者 i 的一个策略。列出所有可能的策略组合,就得到了策略组合集 S : $S \equiv \times_{i \in N} S_i$ 。一个策略组合是一个向量 $s = (s_1, \dots, s_i, \dots, s_n) \in S$ 。集合 $S_{-i} \equiv \times_{j \in N \setminus \{i\}} S_j$ 由参与者 i 以外的所有参与者的策略集构成。这一集合的元素 s_{-i} , 是参与者 $N \setminus \{i\}$ 的一个策略组合。我们常把 s 表示为 (s_i, s_{-i}) 。在后文中,我们将发展策略的定义,允许参与者在纯策略集上做出随机选择。

(3) 收益:标准式博弈中用到的效用函数是定义在 S 上的概率分布集上的 von Neumann-Morgenstern 效用函数。每个参与者都有一个定义在策略组合集上的效用函数: $u_i: S \rightarrow \mathbb{R}$ 。我们有时把 i 的效用函数写为 $u_i(s_i, s_{-i})$ 。根据第 3 章中的定义, $u_i(\cdot)$ 是贝努里效用函数。给定 S 上的某个概率分布,参与者能算出在这一概率分布下自己得到的期望效用。在标准式博弈中,收益可以被理解为期望效用。例如,与两个执法机构都在抓捕核心分子的策略相对应的 2 个单位的收益,就可以被理解为从抓捕恐怖分子的概率分布中得到的期望效用——只是在这一概率分布中,两个机构都在抓捕核心分子。

对标准式博弈,我们也可以这样理解:在阶段 1,每个参与者 $i \in N$ 选择自己的策略 $s_i \in S_i$; 阶段 2,每个参与者 i 得到收益 $u_i(s)$, 其中 $s = (s_1, \dots, s_n)$ 。在后面我们将看到,多阶段博弈可以被表示为标准式博弈。

有了上述定义,我们就可以把标准式博弈写为 $\langle N, \{S_i, u_i(\cdot, \dots, \cdot)\}_{i \in N} \rangle$, 或简写为 $\langle N, S, u \rangle$, 其中,不加下标的 u 表示效用函数向量 $(u_1(\cdot), \dots, u_n(\cdot))$ 。仅

涉及两个参与者的博弈,可以表示为矩阵形式(就像前面所做的那样)。

为深入理解上述概念,我们用标准式博弈表述前面的两个例子。在囚徒两难博弈中, $N = \{\text{参与者 } 1, \text{参与者 } 2\}$, $S_1 = S_2 = \{\text{不认罪, 认罪}\}$, 收益函数是:

$$u_i(s_i, s_{-i}) = \begin{cases} -8, & \text{如果 } s_i = s_{-i} = \text{认罪} \\ -4, & \text{如果 } s_i = s_{-i} = \text{不认罪} \\ -1, & \text{如果 } s_i = \text{认罪} \quad \text{而} \quad s_{-i} = \text{不认罪} \\ -25, & \text{如果 } s_i = \text{不认罪} \quad \text{而} \quad s_{-i} = \text{认罪} \end{cases}$$

在抓捕恐怖分子博弈中, $N = \{\text{CIA}, \text{FBI}\}$, $S_1 = S_2 = \{\text{抓捕核心分子, 抓捕行动人员}\}$,

$$u_i(s_i, s_{-i}) = \begin{cases} 2, & \text{如果 } s_i = s_{-i} = \text{抓捕核心分子} \\ 1, & \text{如果 } s_i = s_{-i} = \text{抓捕行动人员} \\ 1, & \text{如果 } s_i = \text{抓捕行动人员} \quad \text{而} \quad s_{-i} = \text{抓捕核心分子} \\ 0, & \text{如果 } s_i = \text{抓捕核心分子} \quad \text{而} \quad s_{-i} = \text{抓捕行动人员} \end{cases}$$

我们还可以用矩阵表示这两个标准式博弈。对于只有两人参与的博弈,如果策略空间为有限集,则标准式表述形式和博弈矩阵表述形式之间的关系,可以概括为表 5.3 中的模式,其中 $N = \{1, 2\}$, $S_1 = \{s_{11}, \dots, s_{1l}\}$, $S_2 = \{s_{21}, \dots, s_{2k}\}$ 。

表 5.3 一般标准式博弈

$1 \setminus 2$	s_{21}	s_{22}	...	s_{2k}
s_{11}	$u(s_{11}, s_{21})$	$u(s_{11}, s_{22})$	...	$u(s_{11}, s_{2k})$
s_{12}	$u(s_{12}, s_{21})$	$u(s_{12}, s_{22})$	...	$u(s_{12}, s_{2k})$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
s_{1l}	$u(s_{1l}, s_{21})$	$u(s_{1l}, s_{22})$	...	$u(s_{1l}, s_{2k})$

在参与者多于两人时,就很难用矩阵表示博弈,因为我们无法用二维矩阵表示策略组合。但是,我们有时依然有办法做到。设在抓捕恐怖分子的博弈中,美国国家安全局(NSA)也参与进来,它的策略集与另两个机构相同,即 $s_3 = \{\text{抓捕核心分子, 抓捕行动人员}\}$ 。抓捕核心分子需要至少两个机构合作,但各个机构依然能够独立抓到行动人员。FBI 的收益函数为:

$$u_i(s_i, s_{-i}) = \begin{cases} 2, & \text{如果 } s_2 = \text{抓捕核心分子} \quad \text{或} \quad s_3 = \text{抓捕核心分子} \\ 1, & \text{如果 } s_1 = \text{抓捕行动人员} \\ 0, & \text{如果 } s_1 = \text{抓捕行动人员} \quad \text{且} \quad s_3 = \text{抓捕行动人员} \end{cases}$$

这一博弈用表 5.4 和表 5.5 中的两个矩阵表示。

表 5.4 如果 NSA 抓捕核心分子

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	2, 2, 2	2, 1, 2
抓捕行动人员	1, 2, 2	1, 1, 0

在表 5.4 中给出的收益向量所对应的策略组合中,NSA 抓捕核心分子;在表 5.5 中给出的收益向量所对应的策略组合中,NSA 抓捕行动人员。在策略空间为有限集的三人标准式博弈,我们可以用多个收益矩阵将其表示出来,每个收益矩阵对应参与者 1 的一个可能的策略。^①

表 5.5 如果 NSA 抓捕行动人员

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	2, 2, 1	0, 1, 1
抓捕行动人员	1, 0, 1	1, 1, 1

到目前为止,在我们所讨论的博弈中,不存在不确定性。在上述每个博弈中,参与者都知道策略集和收益,知道对手知道博弈的这些构成要素,知道对手知道他们知道,知道对手知道他们知道他们知道,如此等等。这种无限循环被称为共同知识(common knowledge)。标准式博弈还适用于更复杂的环境——参与者不知道各个策略组合相对应的收益。但是,他们知道各个策略组合相对应的期望效用,并且知道他们知道这一期望效用,如此等等。我们来看另一种抓捕恐怖分子模型。在模型中,参与者知道能否抓到恐怖分子是带有一些不确定性的。在这里,两个参与者都对给定策略下抓到恐怖分子的概率有着自己的推断。我们用表 5.6 表示此时的博弈矩阵,表中的收益为期望效用。

在表 5.6 中,10 是抓到核心分子时带来的效用;如果两个机构合作抓捕核心分子,抓到的概率是 $\frac{1}{5}$ 。0 是没有抓到核心分子时所产生的收益。当两个机构合作抓捕核心分子时,无功而返的概率是 $\frac{4}{5}$ 。矩阵中的收益数字和前面矩阵中的数字相同,但是,这个矩阵明确指出了收益如何决定于两个机构的策略和随机事件的。

表 5.6 不确定性条件下抓捕恐怖分子

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	$10 \times \frac{1}{5} + 0 \times \frac{4}{5}, 10 \times \frac{1}{5} + 0 \times \frac{4}{5}$	$0, 6 \times \frac{1}{6} + 0 \times \frac{5}{6}$
抓捕行动人员	$6 \times \frac{1}{6} + 0 \times \frac{5}{6}, 0$	$6 \times \frac{1}{6} + 0 \times \frac{5}{6}, 6 \times \frac{1}{6} + 0 \times \frac{5}{6}$

① 对参与者数量超过三人的博弈,更难用矩阵表示,这需要使用更多的页面和更多的衍射图。

在这个例子中,所有参与者对随机事件不相同的概率推断。所以,它依然是对称信息博弈。在许多策略性环境中,对于各种结果的发生概率,参与者可能有不同的概率推断和信息。下一章将讨论这种不完全信息博弈。

5.2 标准式博弈的解

博弈论模型的目的之一是预测参与者所选择的 S 中的元素。它的另一个目标是,排除 S 中与理性不一致的元素。我们说过,在囚徒两难博弈中, {认罪,认罪} 是可能的结果;在抓捕恐怖分子博弈中, {抓捕核心分子,抓捕核心分子} 或 {抓捕行动人员,抓捕行动人员} 是可能的结果。我们现在探讨这些预测所依托的一般原理。

5.2.1 剔除劣策略

理性行为的第一个特征是对对手的所有可能的策略,某参与者的某个策略,如果比另一策略带来更高的收益,他就不应该选择这后一策略。在表 5.1 中的囚徒两难博弈中,我们的解的前提是,参与者 1 永远不会使用不认罪这一策略,因为无论参与者 2 选择哪一策略,不认罪带给参与者 1 的效用严格低于认罪所实现的效用。就是说,对参与者 1,不认罪策略严格劣于认罪策略。同样地,“不认罪”策略对参与者 2 也是严格劣的。不包含严格劣策略的唯一策略组合是 {认罪,认罪}。下面给出了严格占优的定义。

定义 5.1 (纯策略中的严格占优)对参与者 i ,策略 s_i 严格劣于策略 s'_i 当且仅当对所有的 $s_{-i} \in S_{-i}$, $u_i(s_i, s_{-i}) < u_i(s'_i, s_{-i})$ 。

我们来定义严格劣策略已经被剔除掉所剩下的策略组合。

定义 5.2 (纯策略中通过严格占优进行的剔除)给定策略组合 $s = (s_i, s_{-i})$,如果对任意的 $i \in N$, s_i 不是严格劣的,则说这一策略组合与通过严格占优进行的剔除相一致。

一般情况下,我们不可能仅通过剔除劣策略得到唯一解。但是,通过反复剔除劣策略,我们能够提高预测的准确性。在反复剔除严格劣策略的方法中,我们要对在此前的剔除严格劣策略过程中生存下来的策略,继续剔除严格劣策略。为说明这种方法,我们来看表 5.7 中的博弈矩阵。

表 5.7

1\2	左	中	右
上	1, 0	1, 2	0, 1
下	0, 3	0, 1	2, 0

在这一博弈中,参与者 1 没有严格劣策略。

对参与者 2,策略右劣于策略中,因为在对手使用策略上时,策略中带给参与者 2 两个单位收益而策略右只带来 1 个单位的效用;在对手使用策略下时,策略中和右为参与者 2 产生的收益分别是 1 和 0。所以,参与者 2 不会使用策略右。如

果参与者 1 认识到参与者 2 不会选择策略右,他就会意识到现在的博弈为表 5.8

中的博弈。

在表 5.8 中的博弈中,对参与者 1,策略下劣于策略上(无论对手使用哪种策略,上和下带给参与者 1 的收益分别是 1 和 0)。参与者 2 知道参与者 1 会选择策略上,他就会选择策略中。 $\{\text{上}, \text{中}\}$ 是与反复剔除严格劣策略相一致的唯一的解。下面的定义给出了这一过程的正式定义。

表 5.8

1\2	左	中
上	1, 0	1, 2
下	0, 3	0, 1

定义 5.3 (反复剔除严格劣策略)对标准式博弈 $\Gamma^0 = \langle N, S^0, u^0 \rangle$, 下面的算法反复剔除严格劣策略。

在第 t 个阶段,任意选择参与者 $i' \in N \setminus \{i^{t-1}\}$, 从 S_i^{t-1} 中剔除在博弈 Γ^{t-1} 中严格劣的每个策略。将剔除后生存下来的策略集记为 S_i^t 。对 $j \in N \setminus \{i'\}$, 定义 $S_j^t \equiv S_j^{t-1}$; 对每个 $z \in N$, 定义 u_z^t 为 u_z^0 在 S' 上的限制。

在阶段 t , 如果每个 $i' \in N$ 都没有在博弈 Γ^{t-1} 中严格劣的策略, 则称集合 S^{t-1} 为在反复剔除严格劣策略中生存下来的结果集。

无论参与者参与剔除的先后顺序如何, 最终都得到同一个非劣策略集。对这一过程的行为解释是, 参与者都以下面的方式推理:

我知道对手不使用严格劣策略, 并且我知道对手知道我不使用严格劣策略。因此, 我们都从前 n 次剔除中存活下来的策略空间中做出选择。我知道对手不使用在经历过 n 次剔除后形成的这一博弈中的严格劣的策略, 并且我知道对手知道我不使用在这一博弈中严格劣的策略……如此等等。

我们还可以用弱占优概念找出博弈结果, 但这一概念需要更强的前提条件。我们说某个参与者有一个弱劣策略, 如果他还有另一个策略, 对对手的所有的策略组合, 后一策略能产生至少和弱劣策略相同的收益, 并且, 在对手使用某一策略组合时, 后一策略产生比弱策略更高的收益。

定义 5.4 (纯策略上的弱占优)策略 s_i 弱劣于策略 s'_i , 当且仅当对所有的 $s_{-i} \in S_{-i}$, $u_i(s_i, s_{-i}) \leq u_i(s'_i, s_{-i})$, 并且至少对一个 $s_{-i} \in S_{-i}$, $u_i(s_i, s_{-i}) < u_i(s'_i, s_{-i})$ 。

通过弱占优概念进行剔除, 类似通过严格占优概念进行剔除。通过弱占优策略进行剔除的方法, 在政治科学中的一个重要应用, 出现在多数通过的投票博弈中。设有 n 个人 (n 为奇数) 对候选人 D 和 R 投票。每个人, 在他所偏好的候选人胜出时, 得到收益 1; 否则, 得到收益 0。设策略 $s_i = 1$ 表示投票给 D , $s_i = 0$ 表示投票给 R 。由于多数通过规则决定胜出者, 所以, 偏好 D 的投票者得到的收益是:

$$\text{在 } \sum s_i > \frac{(n+1)}{2} \text{ 时, } u_D = 1;$$

在其他情况中, $u_D = 0$

而偏好 R 的投票者的收益是 $1 - u_D$ 。

在这一博弈中,没有严格劣策略。除非有 $\frac{n-1}{2}$ 个人选择 $s_i = 1$, 有 $\frac{n-1}{2}$ 个人选择 $s_i = 0$, 否则,任何人的效用都不受他自己的选择影响。所以,在几乎所有的策略组合下,投票人没有严格偏好。但是,在 S_{-i} 中产生平局的策略组合上,第 i 个投票人有严格偏好——他会投票给其胜出能为他带来最大效用的候选人。也就是说,投票给自己所偏好的候选人,弱优于投票给自己所不喜欢的候选人。这一弱占优策略,只在某一个策略组合上(即在对手的选择导致平局时)产生更大效用,而在所有其他策略组合上,产生相同的效用。如果剔除掉弱劣策略,则每个人都投票给他所偏好的候选人,被多数人所偏好的候选人胜出。

建立在占优概念上的解有其诱人之处,但是,这种方法有时无法剔除掉任何策略。前面所有形式的抓捕恐怖分子博弈,都不包含严格劣或弱劣策略。如果仅使用占优概念,则其中所有的策略组合都是可行的。

5.2.2 Nash 均衡

John Nash 为博弈论做出的基础性贡献是为标准式博弈提供了一种解概念,这一概念有广泛的应用。Nash 的解是一个策略组合 s^* , 对所有的 $i \in N$, 参与者 i 的策略 s_i^* 是对对手的策略 s_{-i}^* 的“最优反应”。因此,在博弈论中,一个最重要的概念便是最优反应对应。^①

定义 5.5 参与者 $i \in N$ 的最优反应对应是映射 $b_i: S_{-i} \rightarrow S_i$: 对每个 $s_{-i} \in S_{-i}$, $b_i(s_{-i}) = \{s_i \in S_i: \text{对每个 } s'_i \in S_i, u_i(s_i, s_{-i}) \geq u_i(s'_i, s_{-i})\}$ 。

对对手的组合 s_{-i} 的最优反应,是在对手使用 s_{-i} 时,最大化参与者 i 的效用的所有策略构成的集合。于是,对每个策略组合 s_{-i} , 最优反应对应给出了参与者 i 的最优反应集。我们来看前述例子中的最优反应对应。在囚徒两难博弈中,最优反应是:

$$b_1(\text{认罪}) = \{\text{认罪}\}$$

$$b_1(\text{不认罪}) = \{\text{认罪}\}$$

$$b_2(\text{认罪}) = \{\text{认罪}\}$$

$$b_2(\text{不认罪}) = \{\text{认罪}\}$$

在两个执法机构抓捕恐怖分子博弈中,最优反应对应是:

$$b_1(\text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$$

$$b_1(\text{抓捕行动人员}) = \{\text{抓捕行动人员}\}$$

^① 书末的数学附录介绍了对应。

$$b_2(\text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$$

$$b_2(\text{抓捕行动人员}) = \{\text{抓捕行动人员}\}$$

在这些例子中,最优反应集中只有一个行动,但是在许多情况中,最优反应可能是一组策略。在多数通过规则下的投票博弈中,除非其他人的策略组合产生平局,否则,投票给任何一位候选人都是最优反应。正式地说,偏好候选人 D 的投票人 i 的最优反应对应为:

$$b_i\left(s_{-i}: \sum_{j \in N \setminus \{i\}} s_j = \frac{n-1}{2}\right) = \{1\}$$

$$b_i\left(s_{-i}: \sum_{j \in N \setminus \{i\}} s_j \neq \frac{n-1}{2}\right) = \{0, 1\}$$

Nash 均衡是个策略组合,在这一组合中,给定对手的策略,每个人都使用其最优反应集中的某个策略。

定义 5.6 在标准式博弈中,(纯策略)Nash 均衡是策略组合 s^* ,对每个 $i \in N$,

$$s_i^* \in b_i(s_{-i}^*)$$

我们再给出 Nash 均衡的另一定义,它并不以最优反应对应为基础。

定义 5.7 标准式博弈的(纯策略)Nash 均衡是策略组合 s^* ,对每个 $s'_i \in S_i$ 和每个 $i \in N$,

$$u_i(s_i^*, s_{-i}^*) \geq u_i(s'_i, s_{-i}^*)$$

这两个定义是等价的。Nash 均衡的概念只是看上去简单。它要求参与者正确地推断出对手的行动,根据这一推断做出最优反应。对 Nash 均衡的另一种解释,则以第二个定义为基础:没有人有动机单方面改变自己的策略从而偏离 Nash 均衡策略组合。

我们应用这两个定义到前面的例子中。

(1) 囚徒两难:对两个参与者来说,认罪策略是给定对手的所有策略,每个人的最优反应集中的唯一的元素。因此,唯一的 Nash 均衡是{认罪,认罪}。

(2) 两机构抓捕恐怖分子博弈:对两个机构来说, $b_1(\text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$,所以,{抓捕核心分子,抓捕核心分子}是 Nash 均衡;并且,对这两个机构, $b_1(\text{抓捕行动人员}) = \{\text{抓捕行动人员}\}$,所以{抓捕行动人员,抓捕行动人员}也是 Nash 均衡。

(3) 在三个执法机构抓捕恐怖分子博弈中,最优反应对应为:

$$b_1(\text{抓捕核心分子}, \text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$$

$$b_1(\text{抓捕行动人员}, \text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$$

$$b_1(\text{抓捕核心分子}, \text{抓捕行动人员}) = \{\text{抓捕核心分子}\}$$

$$b_1(\text{抓捕行动人员}, \text{抓捕行动人员}) = \{\text{抓捕行动人员}\}$$

$b_i(\text{抓捕核心分子}, \text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$ 意味着 $\{\text{抓捕核心分子}, \text{抓捕核心分子}, \text{抓捕核心分子}\}$ 是 Nash 均衡。 $b_i(\text{抓捕行动人员}, \text{抓捕行动人员}) = \{\text{抓捕行动人员}\}$ 意味着 $\{\text{抓捕行动人员}, \text{抓捕行动人员}, \text{抓捕行动人员}\}$ 是 Nash 均衡。 $b_i(\text{抓捕行动人员}, \text{抓捕核心分子}) = \{\text{抓捕核心分子}\}$ 和 $b_i(\text{抓捕核心分子}, \text{抓捕行动人员}) = \{\text{抓捕核心分子}\}$ 则意味着不存在只有两个机构在抓捕核心分子的 Nash 均衡, 虽然只需两个机构相互合作就足以抓到核心分子。

(4) 多数通过规则下的投票博弈: 几乎所有的策略组合都是 Nash 均衡。在所有使得 $\sum_{i \in N} s_i < \frac{n-1}{2}$ 或者使 $\sum_{i \in N} s_i > \frac{n+1}{2}$ 策略组合中, 对所有的 i , $b_i(s) = \{0, 1\}$ 。因此, 每个满足上述条件的策略组合都是 Nash 均衡。还有一种情况是 $\sum_{i \in N} s_i \in \left\{ \frac{n-1}{2}, \frac{n+1}{2} \right\}$ 。假设 $\sum_{j \in N \setminus \{i\}} s_j = \frac{n+1}{2}$ 。满足这一性质的策略组合是 Nash 均衡, 当且仅当所有选择 $s_i = 1$ 的人偏好 D。因为如果不是如此, 就存在某个 i , 他虽选择了 $s_i = 1$ 却偏好 R。但是, 由于 $\sum_{j \in N \setminus \{i\}} s_j = \frac{n-1}{2}$, 所以, $b_i(s_{-i}) = \{0\}$ 。因此, 这种策略组合不是 Nash 均衡。如果 $\sum_{i \in N} s_i = \frac{n-1}{2}$, 则 s 是 Nash 均衡, 当且仅当所有选择 $s_i = 0$ 的人都偏好 R。总之, 在这一博弈中, Nash 均衡集合几乎包含所有策略组合。而在这个集合以外的策略组合上, 一位候选人凭借微弱多数胜出, 有人虽偏好落选的候选人却把票投给了胜出候选人。

我们还要注意 Nash 均衡集, 和在反复剔除严格劣策略中存活下来的策略组合集之间的相似和不同之处。在囚徒两难中, 两种解相同。而在各种形式的抓捕恐怖分子博弈中, Nash 均衡集小于在剔除劣策略中存活下来的策略组合集。在多数通过规则下的投票博弈中, Nash 均衡集小于在剔除严格劣策略中存活下来的策略组合集, 但是大于剔除弱劣策略中生存下来的结果集, 在后一方法下, 有唯一一个结果。

下面的定理明确了 Nash 均衡和在反复剔除严格劣策略中存活下来的策略组合之间的关系。

定理 5.1 如果策略组合 $(s_1^*, \dots, s_N^*)$ 是 Nash 均衡, 则其中任何一个策略都不会在反复剔除严格劣策略中被剔除掉。

证明: 假设这一定理不成立。设在这一 Nash 均衡策略组合中 s_i^* 是首先要被剔除的策略。这意味着存在某个 s_i , 对所有尚未被剔除的 s_{-i} , $u_i(s_i^*, s_{-i}) < u_i(s_i, s_{-i})$ 。根据假设, s_{-i}^* 尚未被剔除, 因此, $u_i(s_i^*, s_{-i}^*) < u_i(s_i, s_{-i}^*)$ 。这与 $(s_1^*, \dots, s_N^*)$ 是 Nash 均衡的前提矛盾。□

定理 5.2 如果策略组合 $(s_1^*, \dots, s_N^*)$ 是在反复剔除严格劣策略中唯一存活下来的策略组合, 则它是唯一的纯策略 Nash 均衡。

证明: 唯一性来自前一定理, 因为任何 Nash 均衡肯定在反复剔除中存活下来。我们来证明存活下来的策略组合肯定是 Nash 均衡。设 $(s_1^*, \dots, s_N^*)$ 不是 Nash 均

衡。则存在 s'_i 使得 $u_i(s'_i, s_{-i}^*) < u_i(s_i^*, s_{-i}^*)$, 但是, s'_i 肯定要被剔除掉。设 s_i 是使 s'_i 被剔除掉的策略。如果 $s_i = s_i^*$, 则存在矛盾。设 $s_i \neq s_i^*$ 。由于 s_i 使 s'_i 被剔除掉, 所以 $u_i(s'_i, s_{-i}^*) < u_i(s_i, s_{-i}^*)$ 。这一过程持续有限次, 就能得到 $u_i(s_i^*, s_{-i}^*) > u_i(s'_i, s_{-i}^*)$ 。□

在前述所有例子中, 都存在至少一个 Nash 均衡。但是, 还有一种可能性: 所有的策略组合都不满足 Nash 均衡条件。分析下面的博弈。假设有两支军队: 一支是防御的军队(D), 另一支是进攻的军队(I)。进攻军队要决定是穿过高山(M)还是穿过平原(P)发起进攻。面对侵略, 防守军队得决定是加强山区防御(M)还是在平原防御(P)。进攻军队如果进攻了没有防御的区域, 得到收益 1; 如果进攻了有防御的地区, 得到收益 -1。防守部队如果正确预测到了对方的进攻方向, 防守部队得到收益 1, 否则得到收益 -1。这一标准式博弈被称为 Colonel Blotto 博弈, 我们将其表示为表 5.9 中的矩阵形式。

表 5.9 Colonel Blotto

D\I	M	P
M	1, -1	-1, 1
P	-1, 1	1, -1

最优反应对应是: $b_D(M) = \{M\}$ 、 $b_D(P) = \{P\}$ 、 $b_I(M) = \{P\}$ 和 $b_I(P) = \{M\}$ 。Nash 均衡要求存在某个策略组合 (s_D, s_I) 使得 $b_I(s_D) = s_I$ 和 $b_D(s_I) = s_D$ 。这一博弈中的最优反应对应并不满足这些条件。给定任何一个对策略, 每个人都有动机选择另一策略。没有 Nash 均衡(或由占优施加的限制), 我们无法预测这一博弈如何进行。

在应用分析中, 我们一般定义一个博弈, 进而阐述其 Nash 均衡集的特征。在分析 Nash 均衡集时, 两个最重要的性质是存在性和唯一性。对应用研究者而言, 有唯一的 Nash 均衡可谓是最理想的情况, 因为这意味着在经验实证预测上没有任何含糊之处。相比之下, 多均衡就不是很理想的情况, 因为模型产生的预测并不明确。而没有均衡的情况最糟糕, 因为此时的模型无法形成任何预测。在后面我们将看到, 任何博弈, 只要策略数量和参与者数量有限, 如 Colonel Blotto 博弈那样, 参与者如能在策略集上随机做出选择, 就足以保证 Nash 均衡的存在。

5.3 应用: Hotelling 政治竞争模型

Hotelling(1929)政治竞争模型, 由 Downs(1957)提出, 是政治科学中应用最广泛的博弈模型之一。一个小镇要决定在何处建所学校。镇里的村民沿着一条一英里长的马路均匀分布, 每个人都希望学校离自己的家尽可能地近。因此, 投票人的理想校址在 $[0, 1]$ 上均匀分布。小镇举行选举, 有两位候选人竞选镇长职位, 他们的竞选变量便是学校的位置。胜出的候选人将在他所承诺的地点建学校, 自己从中得到 1 单位的收益。落选候选人得到的收益为 -1。一旦出现平局, 则通过抛硬币决定胜出者, 也就是说, 出现平局时每个候选人胜出的概率都为 0.5。为使讨论尽可能地简单, 我们假定在这一博弈中, 投票人没有策略性行为, 而始终投票给

所承诺的校址离自己家最近的那位候选人。^①

两位候选人的策略集是 $S_1 = S_2 = [0, 1]$ 。假设投票人投票给承诺最近地点的候选人,在此前提下,对任意的策略组合 (s_1, s_2) , 我们计算两个候选人的得票数。由于投票人均匀分布,所以在任何区间中,投票人的数量等于区间长度。因此,如果 $s_2 > s_1$, 在点 $\frac{s_1 + s_2}{2}$ 左边的所有投票人都投票给候选人 1, 他的得票数是 $\frac{s_1 + s_2}{2}$ 。另外的 $1 - \frac{s_1 + s_2}{2}$ 个投票人支持候选人 2。如果 $s_1 > s_2$, 则候选人 2 的得票数是 $\frac{s_1 + s_2}{2}$, 候选人 1 得到剩余票数。对任意的策略组合, 候选人的收益是:

$$u_1(s_1, s_2) = \begin{cases} 1, & \text{如果 } s_1 < s_2 \text{ 且 } \frac{s_1 + s_2}{2} > 0.5, \text{ 或者 } s_1 > s_2 \text{ 且 } \frac{s_1 + s_2}{2} < 0.5 \\ 0, & \text{如果 } s_1 = s_2 \\ -1, & \text{如果 } s_1 < s_2 \text{ 且 } \frac{s_1 + s_2}{2} < 0.5, \text{ 或者 } s_1 > s_2 \text{ 且 } \frac{s_1 + s_2}{2} > 0.5 \end{cases}$$

$$u_2(s_1, s_2) = -u_1(s_1, s_2)$$

这一博弈的唯一的 Nash 均衡是 $s_1 = s_2 = 0.5$ 。要说明这一点, 我们首先计算候选人 1 的最优反应对应。设 $s_2 < 0.5$, 候选人 1 所选择的地点只要能产生 0.5 或更多的票数, 就肯定胜出。就是说, $b_1(s_2) = (s_2, 1 - s_2)$ 。设 $s_2 > 0.5$, 类似的计算表明, $b_1(s_2) = (1 - s_2, s_2)$ 。设 $s_2 = 0.5$, 则对候选人 1 来说, 最好的结果是选择 0.5, 即 $b_1(0.5) = 0.5$, 从而产生平局。候选人 2 的情况与候选人 1 对称, 所以他的最优反应对应是:

$$b_2(s_1) = (s_1, 1 - s_1), \text{ 如果 } s_1 < 0.5$$

$$b_2(s_1) = (1 - s_1, s_1), \text{ 如果 } s_1 > 0.5$$

$$b_2(s_1) = s_1, \text{ 如果 } s_1 = 0.5$$

由于 $b_1(0.5) = 0.5$ 且 $b_2(0.5) = 0.5$, 所以 $s_1^* = s_2^* = 0.5$ 是 Nash 均衡。这里还要证明它是唯一的 Nash 均衡。设 $s_1^* = s_2^* \neq 0.5$ 。这一策略组合不是 Nash 均衡, 因为 $b_1(s_2^*)$ 中不包含 s_2^* 。设 $s_1^* < s_2^*$, 则 $b_1(s_2^*) = (1 - s_2^*, s_2^*)$ 而 $b_1(s_1^*) = (s_1^*, 1 - s_1^*)$ 。综合这两个条件, Nash 均衡要求 $1 - s_2^* < s_1^*$ 且 $s_2^* < 1 - s_1^*$ 。但这两个不等式相互矛盾, 因为它们分别要求 $s_2^* > 1 - s_1^*$ 和 $s_2^* < 1 - s_1^*$ 。设 $s_1^* > s_2^*$, 此时的结论与前一种情况相同。因此, $s_1^* = s_2^* = 0.5$ 为唯一的 Nash 均衡。

利用 Nash 均衡的第二个定义(相对于 Nash 均衡, 没有人严格偏好非均衡策略)能更直观地证明这一结论。很明显, 没有候选人会偏离 $s_1^* = s_2^* = 0.5$, 因为这

① 这一假定事实上并不重要; 投票给承诺了更近地点的候选人, 与在投票人在了解候选人所承诺的地点后再决定如何投票的博弈中出现的均衡相一致。

种偏离使他的收益从 0 下降到 -1。设 (s_1^*, s_2^*) 为可能的均衡。在这一均衡中如果有人胜出,另一位候选人就可以选择 0.5——对他而言,这样做的最糟糕的结果不过是双方打成平局。因此,在任何均衡中,结果肯定都是平局。而且,除非 $s_1^* = s_2^* = 0.5$, 否则任何人都可以通过选择 0.5 而胜出。策略组合 $s_1^* = s_2^* = 0.5$ 是唯一的 Nash 均衡。

剔除弱劣策略也能产生相同的结果。一位候选人,如果选择 0.5,最差的结果是出现平局。当对手也选择 0.5 时,就会出现平局。其他任何策略都将导致落选。因此, $s_i = 0.5$ 弱优于其他所有策略。

5.3.1 候选人最大化选票数

假定每个候选人不是最大化自己的胜选概率,而最大化自己的得票数。现在的收益函数是:

$$u_1(s_1, s_2) = \begin{cases} \frac{s_1 + s_2}{2}, & \text{如果 } s_1 < s_2 \\ 0.5, & \text{如果 } s_1 = s_2 \\ 1 - \frac{s_1 + s_2}{2}, & \text{如果 } s_1 > s_2 \end{cases}$$

$$u_2(s_1, s_2) = -u_1(s_1, s_2)$$

策略组合 $s_1^* = s_2^* = 0.5$ 依然是唯一的 Nash 均衡。我们从最优反应对应着手分析。设 $s_2 < 0.5$ 。此时,候选人 1 偏好大于 s_2 的最小位置。由于策略集是连续统,所以不存在这样的位置。设 $s_2 > 0.5$ 。此时,候选人 1 会选择小于 s_2 的最大位置,但这样的位置并不存在。候选人 2 身临同样的处境。因此, $b_i(s_{-i} \neq 0.5) = \emptyset$, 即最优反应对应为空集。设 $s_2 = 0.5$ 。此时,候选人 1 如果选择 0.5,就得到 0.5 的票数,从而产生平局;如果他没有选择 0.5,他的得票数就更低。所以,0.5 是候选人 1 的最优反应。候选人 2 面对着同样的激励。所以, $b_i(s_{-i} = 0.5) = 0.5$ 。 $s_1^* = s_2^* = 0.5$ 显然是 Nash 均衡。由于对其他任何策略组合而言,最优反应集都是空集,所以, $s_1^* = s_2^* = 0.5$ 是唯一的 Nash 均衡。

5.3.2 极端候选人

我们看另一种形式的 Hotelling 模型。在这种模型中,候选人是极端主义者,因为他们关心胜出的候选人执行何种政策。候选人 1 希望学校尽可能地靠近点 0,他从位置 x 中获得的效用是 $-x^2$ 。候选人 2 希望结果尽可能地靠近点 1,他的效用函数为 $-(1-x)^2$ 。和前面一样,候选人所选择的竞选承诺是在其当选后必须兑现的承诺。^①

① 候选人在当选后可能会选择自己的理想点,而置竞选承诺于不顾。这里不考虑这种可能性。

在选举出现平局时,通过抛硬币决定胜出者,他兑现自己的承诺。由于候选人都希望政策趋于端点位置,所以,最终结果似乎不在中间位置上。实际情况并非如此, $s_1^* = s_2^* = 0.5$ 依然是唯一的 Nash 均衡。

两个候选人的收益函数分别为:

$$u_1(s_1, s_2) = \begin{cases} -s_1^2, & \text{如果 } s_1 < s_2 \text{ 且 } \frac{s_1 + s_2}{2} > 0.5 \text{ 或者 } s_1 > s_2 \text{ 且 } \frac{s_1 + s_2}{2} < 0.5 \\ -0.5s_1^2 - 0.5s_2^2, & \text{如果 } s_1 = s_2 \\ -s_2^2, & \text{如果 } s_1 < s_2 \text{ 且 } \frac{s_1 + s_2}{2} < 0.5 \text{ 或者 } s_1 > s_2 \text{ 且 } \frac{s_1 + s_2}{2} > 0.5 \end{cases}$$

$$u_2(s_1, s_2) = \begin{cases} -(1-s_1)^2, & \text{如果 } s_1 < s_2 \text{ 且 } \frac{s_1 + s_2}{2} > 0.5 \text{ 或者 } s_1 > s_2 \text{ 且 } \frac{s_1 + s_2}{2} < 0.5 \\ -0.5(1-s_1)^2 - 0.5(1-s_2)^2, & \text{如果 } s_1 = s_2 \\ -(1-s_2)^2, & \text{如果 } s_1 < s_2 \text{ 且 } \frac{s_1 + s_2}{2} < 0.5 \text{ 或者 } s_1 > s_2 \text{ 且 } \frac{s_1 + s_2}{2} > 0.5 \end{cases}$$

首先分析候选人 1 的最优反应。如果 $s_2 < 0.5$, 小于 s_2 的任何竞选承诺都无法击败 s_2 。所以,候选人 1 的最优反应是选择 s_2 或者选择输给 s_2 的某个方案,即 $b_1(s_2 < 0.5) = [0, s_2] \cup (1-s_2, 1]$ 。如果 $s_2 > 0.5$, 候选人 1 会选择能击败 s_2 的某个最小的位置。但是,根据上一节中的分析,这样的位置并不存在,所以, $b_1(s_2 > 0.5) = \emptyset$ 。另一方面,如果候选人 2 选择 $s_2 > 0.5$, 则对击败 s_2 的任何政策 s_1 而言,这样的 s_2 对候选人 2 是次优的。同样的道理, $b_2(s_1 > 0.5) = [s_1, 1] \cup [0, 1-s_1)$ 且 $b_2(s_1 < 0.5) = \emptyset$ 。如果候选人 1 选择 $s_1 < 0.5$, 则对候选人 2 所选择的击败 s_1 的任何政策而言,这样的 s_1 对候选人 1 是次优的。再看对 $s_i = 0.5$ 的最优反应。此时,对手的任何选择都注定落败,最终的效用为 -0.25 。而如果对手选择 0.5 , 就有 $s_1 = s_2 = 0.5$, 于是要通过抛硬币决定胜负。抛硬币的概率分布带给两位候选人的期望效用都是 -0.25 。所以,最优反应是 $b_i(0.5) = [0, 1]$ 。

根据上述最优反应对应, $s_1^* = s_2^* = 0.5$ 是 Nash 均衡,因为对两个候选人来说,都有 $0.5 \in b_i(0.5)$ 。现在来证明均衡的唯一性。由于 $b_1(s_2 > 0.5) = \emptyset$ 且 $b_2(s_1 < 0.5) = \emptyset$, 所以,要成为 Nash 均衡,唯一可能的是 $s_1^* > 0.5 > s_2^*$ 。这一不等式与 $b_1(s_2 < 0.5) = [0, s_2] \cup (1-s_2, 1]$ 和 $b_2(s_1 > 0.5) = [s_1, 1] \cup [0, 1-s_1)$ 相结合,得到 $s_1^* > 1-s_2^*$ 和 $s_2^* < 1-s_1^*$ 。最后这两个不等式不可能同时成立,所以, $s_1^* = s_2^* = 0.5$ 是唯一可能的 Nash 均衡。

5.3.3 不确定环境中的意识形态候选人

极端候选人也会拉拢中间选民,因为如果没有中间选民的支持,没人能够胜出。现在分析另一种情况。如果两位候选人都不知道哪种竞选方案能得到中间选民的支持,又会出现怎样的结果呢?^①

① 在学习本节前,你可能需要重温数学附录中关于优化的介绍。

现在,两位候选人都不知道中间选民的位置。他们不知道选民在单位区间上是均匀分布的,而认为中间选民的位置是在 $[0, 1]$ 区间上的均匀分布中随机抽取的。这是 Wittman(1977)和 Calvert(1985)所提出的模型。设 $\Omega = [0, 1]$ 表示中间选民可能的位置构成的集合。候选人的偏好是共同知识,所以,博弈中的所有不确定性都反映于 $[0, 1]$ 上的概率分布 $F(\omega) = \omega$ 。一旦我们把候选人的期望效用表示为候选人策略的函数,我们就把这一问题转变为了标准式博弈。假定候选人在政策上的偏好是二次形式, $u_1(x) = -x^2$, $u_2(x) = -(1-x)^2$ 。给定两个人的政策 s_1 和 s_2 , 如果中间选民所在的位置更接近 s_1 而不是 s_2 , 则候选人 1 胜出。所以, 如果 $s_1 < s_2$ 并且中间选民所在的位置小于 $\frac{s_1 + s_2}{2}$, 候选人 1 胜出。因为中间选民的位置均匀分布, 所以候选人 1 胜出为概率是 $\frac{s_1 + s_2}{2}$ 。利用这一事实, 可以把候选人 1 和候选人 2 的期望效用表示为:

$$Eu_1(s_1, s_2) = \begin{cases} -s_1^2 \left(\frac{s_1 + s_2}{2} \right) - s_2^2 \left(1 - \frac{s_1 + s_2}{2} \right), & \text{如果 } s_1 < s_2 \\ -s_2^2 \left(\frac{s_1 + s_2}{2} \right) - s_1^2 \left(1 - \frac{s_1 + s_2}{2} \right), & \text{如果 } s_1 > s_2 \end{cases}$$

$$Eu_2(s_1, s_2) = \begin{cases} -(1-s_1)^2 \left(\frac{s_1 + s_2}{2} \right) - (1-s_2)^2 \left(1 - \frac{s_1 + s_2}{2} \right), & \text{如果 } s_1 < s_2 \\ -(1-s_2)^2 \left(\frac{s_1 + s_2}{2} \right) - (1-s_1)^2 \left(1 - \frac{s_1 + s_2}{2} \right), & \text{如果 } s_1 > s_2 \end{cases}$$

我们求解 Nash 均衡。设候选人 1 知道, 候选人 2 的位置为 $z \geq \frac{1}{2}$ 。他可以通过选择 $s_1 \in [0, z]$ 最大化:

$$\max_{s_1} \left\{ -s_1^2 \left(\frac{s_1 + z}{2} \right) - z^2 \left(1 - \frac{s_1 + z}{2} \right) \right\}$$

我们不必考虑选择 $s_1 > z$ 这种可能性, 因为这种策略总是劣于 $s_1 = z$ 。①为找出最优的 s_1 , 我们把目标函数对 s_1 求导并令导数等于 0。一阶条件为:

$$-\frac{3}{2}s_1^2 - zs_1 + \frac{z^2}{2} = 0$$

求解 s_1 , 得到两个解, 其中只有一个位于区间 $[0, 1]$ 中。这个解给出了最优反应函数:

$$s_1(s_2) = \frac{s_2}{3}$$

要保证这一解产生局部最大值(而不是局部最小值或鞍点), 我们需要验证在

① 实际上, 如果候选人 1 选择了 $s_1 > z$, 他会希望候选人 2 胜出。

这一解上,目标函数的二阶导数的值为负。把目标函数对 s_1 求导,并把 z 看作是常数,得到二阶导数,为 $-2z$ 。对任何的 $z \in (0, 1]$, 二阶导数值是负的。

同样地,把 s_1 看作是固定参数 $z \leq \frac{1}{2}$, 求解候选人 2 的最优选择:

$$\max_{s_2} \left\{ -(1-z)^2 \left(\frac{z+s_2}{2} \right) - (1-s_2)^2 \left(1 - \frac{z+s_2}{2} \right) \right\}$$

对候选人 2 的目标函数求导,可以找到最优的 $s_2 \in [z, 1]$ 。求得的解为:

$$s_2(s_1) = \frac{2}{3} + \frac{1}{3}s_1$$

这里请你验证二阶条件肯定成立。Nash 均衡 (s_1^*, s_2^*) 便是下面的方程组的解:

$$s_1^* = \frac{1}{3}s_2^*$$

$$s_2^* = \frac{2}{3} + \frac{1}{3}s_1^*$$

这个方程组的唯一的解是 $(s_1^* = \frac{1}{4}, s_2^* = \frac{3}{4})$ 。因此,当候选人有极端倾向性且在选民的偏好上存在着不确定性时,两个候选人选择不同的竞选纲领。在确定性模型中,一个候选人的选择如果偏离 0.5,会导致他肯定落选。而在当前的模型中,这种偏离只会导致他的胜出概率下降。极端候选人为使政策纲领接近其理想点,一般愿意在胜出概率上做点牺牲。最终的结果是,两个候选人都根据自己的偏好,使自己的竞选纲领偏离中间选民。由于两个人都是这么做的,所以,每个人的胜出概率依然为 50%。

5.4 Nash 均衡的存在性

在 Colonel Blotto 模型中,纯策略 Nash 均衡未必存在。本节探讨 Nash 均衡的存在性条件。

首先给出存在纯策略均衡的充分条件。第一个概念是效用函数的凹性定义。

定义 5.8 给定函数 $f: X \rightarrow \mathbb{R}$, 其中, X 是凸子集。如果对所有的 $x, y \in X$ 且 $x \neq y$, 对所有的 $\lambda \in (0, 1)$, 和对任意的 $t \in \mathbb{R}$, 如果 $f(x) \geq t$ 且 $f(y) \geq t$, 则 $f(\lambda x + (1-\lambda)y) > t$, 就说 f 是严格拟凹的。

换句话说,一个函数,如果它的上半集是凸的,则它是严格拟凹的。由于严格凸偏好产生单元素最优集(见第 2 章),所以,在博弈中,如果策略空间 S 为凸集且对每个 $s_{-i} \in S_{-i}$, 效用函数 $u_i(s_i, s_{-i})$ 在 s_i 上严格拟凹,则这个博弈中的最优反应对应为函数(即单值对应)。在下一小节中,我们证明下面的结论。

定理 5.3 如果以下条件得到满足,则标准式博弈 (N, S, u) 有 Nash 均衡:

(1) 对每个 $i \in N$, S_i 是欧几里得空间中的凸的和紧的子集;

(2) 对每个 $i \in N$, $u_i: S \rightarrow \mathbb{R}$ 是连续函数;

(3) 对每个 $i \in N$ 和每个 $s'_{-i} \in S_{-i}$, 函数 $u_i(\cdot, s'_{-i}): S_i \rightarrow \mathbb{R}$ 在 S_i 上严格拟凹。

这一定理尽管有用,但是限制性很强,因为许多博弈并不满足其中的假设。例如,Colonel Blotto 博弈不满足其中的条件 1,它的策略集不是凸集(M 和 P 的线性组合不是可用的策略)。在前述各种形式的 Hotelling 博弈中,存在唯一的 Nash 均衡,但是,收益函数并不严格拟凹,因为在很多区域中,收益函数在候选人的可能的竞选方案上是“平的”(落选方案依然是落选方案,因为它偏离了中间选民)。

对策略空间并非凸集的博弈,一种解决办法是允许参与者使用混合策略。所谓的混合策略,我们将其表示为 σ_i ,是纯策略集上的随机化。 $\sigma_i(s_i)$ 表示参与者 i 选择策略 s_i 的概率。参与者 i 的混合策略集,是 S_i 上的概率分布构成的集合,表示为 $\Delta_i = \Delta(S_i)$ 。

混合策略博弈类似纯策略博弈,只是参与者选择的是 $\sigma_i \in \Delta_i$ 而不是 $s_i \in S_i$,并且参与者根据由 σ_i 诱生的概率分布的期望效用比较各个策略。下面是混合策略博弈的正式定义。

定义 5.9 对标准式博弈 $\Gamma = \langle N, S, u \rangle$, 其混合策略博弈为 $\Gamma^m = \langle N, \Delta, u^m \rangle$; 策略集为 $\Delta_i = \Delta(S_i)$ 。对所有的 $i \in N$, $\sigma_i \in \Delta_i$ 表示 i 的一个策略; $\Delta = \times_{i \in N} \Delta_i$; $\sigma_i(s_i)$ 为 σ_i 赋予 s_i 的概率。对所有的 $i \in N$, 期望效用函数 $u_i: \Delta \rightarrow \mathbb{R}$ 为

$$u_i(\sigma_i, \sigma_{-i}) = \sum_{s_{-i} \in S_{-i}} \sum_{s_i \in S_i} u_i(s_i, s_{-i}) \sigma_i(s_i) \sigma_{-i}(s_{-i})$$

标准式博弈所衍生的混合博弈依然属于标准式博弈,所以 Nash 均衡的定义直接适用于混合博弈。

我们用 Colonel Blotto 博弈说明参与者如何使用混合策略。攻防双方现在选择自己的策略集上的概率分布。设 $\sigma_1 = \sigma_1(M)$ 是防御方保护山地的概率, $\sigma_2 = \sigma_2(M)$ 是进攻方进攻山地的概率。下面四个式子给出了各个行动带给双方的期望效用:

$$u_1(M, \sigma_2) = \sigma_2 - (1 - \sigma_2) = 2\sigma_2 - 1$$

$$u_1(P, \sigma_2) = -\sigma_2 + (1 - \sigma_2) = 1 - 2\sigma_2$$

$$u_2(\sigma_1, M) = -\sigma_1 + (1 - \sigma_1) = 1 - 2\sigma_1$$

$$u_2(\sigma_1, P) = \sigma_1 - (1 - \sigma_1) = 2\sigma_1 - 1$$

我们直接计算 $b_1(\sigma_2)$ 和 $b_2(\sigma_1)$ 。如果 $\sigma_2 > \frac{1}{2}$, 防守方会选择防守山地(因为 $u_1(M, \sigma_2) > u_1(P, \sigma_2)$)。如果 $\sigma_2 = \frac{1}{2}$, 防守方对两个纯策略无差异。如果 $\sigma_2 < \frac{1}{2}$, 防守方会防守平原。所以,防守方对进攻方的混合策略 σ_2 的最优反应是:

$$b_1(\sigma_2) = \begin{cases} M, & \text{如果 } \sigma_2 > \frac{1}{2} \\ P, & \text{如果 } \sigma_2 < \frac{1}{2} \\ \{M, P\}, & \text{如果 } \sigma_2 = \frac{1}{2} \end{cases}$$

在 $\sigma_2 = \frac{1}{2}$ 时, $b_1(\sigma_2) = \{M, P\}$, 所以 $\{M, P\}$ 上的任何随机化都构成最优反应。

同样的方法可求得进攻方的最优反应:

$$\text{在 } \sigma_1 < \frac{1}{2} \text{ 时, } u_2(\sigma_1, M) > u_2(\sigma_1, P)$$

$$\text{在 } \sigma_1 > \frac{1}{2} \text{ 时, } u_2(\sigma_1, M) < u_2(\sigma_1, P)$$

$$\text{在 } \sigma_1 = \frac{1}{2} \text{ 时, } u_2(\sigma_1, M) = u_2(\sigma_1, P)$$

所以, 进攻方的最优反应对应是:

$$b_2(\sigma_1) = \begin{cases} M, & \text{如果 } \sigma_1 < \frac{1}{2} \\ P, & \text{如果 } \sigma_1 > \frac{1}{2} \\ \{M, P\} \text{ 或 } \sigma_2, & \text{如果 } \sigma_1 = \frac{1}{2} \end{cases}$$

前文指出, 在参与者使用纯策略时, 这一博弈没有 Nash 均衡。所以, 在构造 Nash 均衡时, 我们只需检验是否存在满足 Nash 条件的 σ_1 和 σ_2 组合。只有当 $\sigma_2 = \frac{1}{2}$ 时, 混合策略 σ_1 才是防守方的最优反应; 只有当 $\sigma_1 = \frac{1}{2}$ 时, 混合策略 σ_2 才是进攻方的最优反应。所以, $\sigma_1 = \frac{1}{2}$ 和 $\sigma_2 = \frac{1}{2}$ 构成了唯一的混合策略 Nash 均衡。

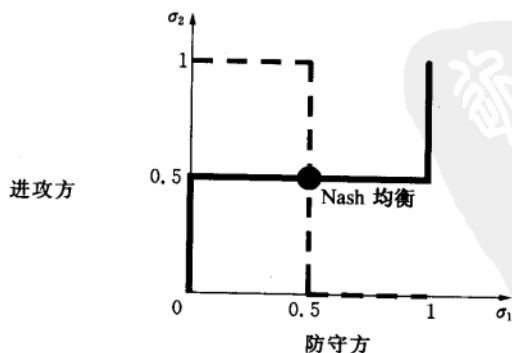


图 5.1 Colonel Blotto 博弈的混合策略 Nash 均衡

在有两个参与者且每个参与者有两种策略的博弈中,画出最优反应函数能帮助找到混合策略均衡。在图 5.1 中,横轴给出了防守方的混合策略,它的取值范围为 $\sigma_1 = 0$ (防守平原)到 $\sigma_1 = 1$ (防守山地)。纵轴给出了进攻方的混合策略,它的取值范围为 $\sigma_2 = 0$ (进攻平原)到 $\sigma_2 = 1$ (进攻山地)。实线是防守方对 σ_2 的最优反应,虚线是进攻方对 σ_1 的最优反应。两条最优反应曲线的唯一交点为 Nash 均衡点: $\sigma_1 = \frac{1}{2}$ 和 $\sigma_2 = \frac{1}{2}$ 。

在有纯策略均衡的博弈中,也可能存在混合策略均衡。在表 5.2 所描述的两执法机构抓捕恐怖分子博弈中,设 σ_1 和 σ_2 分别是 FBI 和 CIA 抓捕核心恐怖分子的概率。各个行动带给两个机构的期望效用分别为:

$$u_1(\text{抓捕核心分子}, \sigma_2) = 2\sigma_2$$

$$u_1(\text{抓捕行动人员}, \sigma_2) = 1$$

$$u_2(\sigma_1, \text{抓捕核心分子}) = 2\sigma_1$$

$$u_2(\sigma_1, \text{抓捕行动人员}) = 1$$

和前面一样,我们通过比较这些效用求得最优反应函数:

$$b_1(\sigma_2) = \begin{cases} \text{抓捕核心分子, 如果 } \sigma_2 > \frac{1}{2} \\ \text{抓捕行动人员, 如果 } \sigma_2 < \frac{1}{2} \\ \sigma_1 \in [0, 1], \text{ 如果 } \sigma_2 = \frac{1}{2} \end{cases}$$

$$b_2(\sigma_1) = \begin{cases} \text{抓捕核心分子, 如果 } \sigma_1 > \frac{1}{2} \\ \text{抓捕行动人员, 如果 } \sigma_1 < \frac{1}{2} \\ \sigma_2 \in [0, 1], \text{ 如果 } \sigma_1 = \frac{1}{2} \end{cases}$$

图 5.2 画出了这些最优反应函数。两条最优反应曲线有三个交点: $\{0, 0\}$ 、 $\{1, 1\}$ 和 $\left\{\frac{1}{2}, \frac{1}{2}\right\}$ 。前两个交点为前面算得的纯策略均衡,第三个交点为混合策略均衡。

混合策略均衡的特点是,每个参与者使用某一策略的概率不是她自己的偏好的函数而是其对手的偏好的函数。这是因为,只有在对自己的各个纯策略无差异时,她才会使用混合策略。于是,其对手通过选择自己的混合策略以确保她是无差异的。因此,对手的偏好决定她的混合策略。这一事实有时产生与直觉相悖的结论。我们修改抓捕恐怖分子博弈。在新的博弈中,抓捕核心分子带给 CIA 的收益高于带给 FBI 的收益。表 5.10 中给出了新的收益数字。

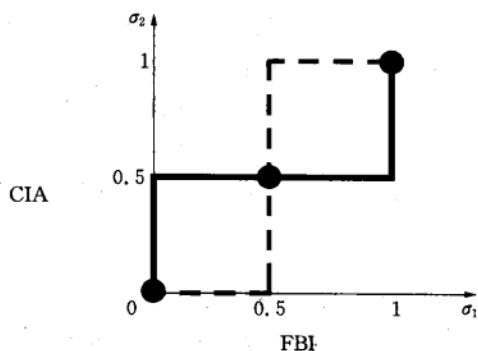


图 5.2 抓捕恐怖分子博弈中的混合策略 Nash 均衡

表 5.10 修改后的抓捕恐怖分子博弈

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	2, 4	0, 1
抓捕行动人员	1, 0	1, 1

在简单的决策理论上形成的直觉是,由于在新博弈中,CIA 从抓捕核心分子行动上获得的收益更高,所以,它更可能去抓捕核心分子。但是,这种认识是错误的。最优反应函数现在是:

$$b_1(\sigma_2) = \begin{cases} \text{抓捕核心分子, 如果 } \sigma_2 > \frac{1}{2} \\ \text{抓捕行动人员, 如果 } \sigma_2 < \frac{1}{2} \\ \sigma_1 \in [0, 1], \text{ 如果 } \sigma_2 = \frac{1}{2} \end{cases}$$

$$b_2(\sigma_1) = \begin{cases} \text{抓捕核心分子, 如果 } \sigma_1 > \frac{1}{4} \\ \text{抓捕行动人员, 如果 } \sigma_1 < \frac{1}{4} \\ \sigma_2 \in [0, 1], \text{ 如果 } \sigma_1 = \frac{1}{4} \end{cases}$$

图 5.3 画出了这些反应。混合策略均衡现在是 $\left\{\frac{1}{4}, \frac{1}{2}\right\}$ 。CIA 偏好上的变化并没有使它以更大的概率抓捕核心分子——它抓捕核心分子的概率依然是 $\frac{1}{2}$ 。但这一变化降低了 FBI 抓捕核心分子的概率。在混合策略均衡中,核心分子被抓住的概率下降了。这一结果背后的逻辑不难理解。在混合策略均衡中,CIA 通过选择 σ_2 ,使 FBI 在抓捕核心分子和抓捕行动人员之间无差异。FBI 的偏好没有改变,所以 σ_2 也没有改变。由于 FBI 通过选择 σ_1 而使 CIA 在自己的两种策略之间无差

异,所以,抓捕核心分子带给 CIA 的效用上升的事实,意味着 FBI 必须降低抓捕核心分子的概率,借此保持这种无差异性。

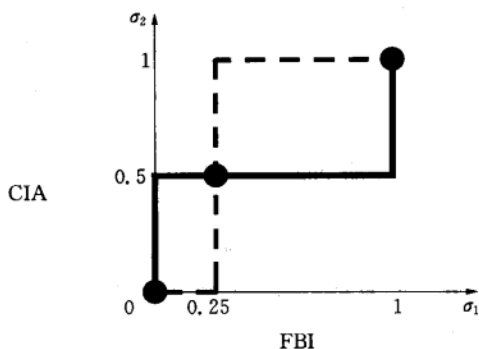


图 5.3 抓捕恐怖分子博弈中的混合策略 Nash 均衡

混合策略均衡尽管有这些不合理的特点,但是在策略集为有限集的博弈中,均衡肯定存在。这就是著名的 Nash 定理:

定理 5.4 (Nash 定理) 给定标准式博弈 $\Gamma = \langle N, S, u \rangle$, 其中 S 是有限集, 其混合形式 $\Gamma^m = \langle N, \Delta, u^m \rangle$ 有至少一个 Nash 均衡。换句话说, 任何有限博弈都有混合策略 Nash 均衡。

这一定理的证明需要一些高深数学概念, 我们把它放到后面证明。我们想指出的是, 每个读者都应该知道这个结论, 但是在初学时, 可以跳过详细的证明部分。

5.5 占优和混合策略

我们推广占优的定义使其能适用于混合策略。

定义 5.10 给定纯策略 $s_i \in S_i$, 如果存在 $\sigma'_i \in \Delta(S_i)$ 使得对每个 $s_{-i} \in S_{-i}$:

$$u_i(\sigma'_i, s_{-i}) > u_i(s_i, s_{-i})$$

则 s_i 是严格劣的。如果存在 σ'_i , 使得对每个 $s_{-i} \in S_{-i}$ 上述式子中弱不等关系成立, 并且存在某个 $s'_{-i} \in S_{-i}$ 使得上式中严格不等关系成立, 则策略 s_i 是弱劣的。^①

这一定义的推广很直接。现在探讨某个策略是否劣于某个混合策略。从下面的例子中, 可以看到, 混合策略占优比纯策略占优更有力, 因为一些纯策略并不劣于其他纯策略, 但劣于混合策略。^②我们来看表 5.11 中的博弈。

① $u_i(\sigma'_i, s_{-i}) = \sum_{s'_i \in S_i} u_i(s'_i, s_{-i}) \sigma'_i(s'_i)$ 。

② 逆命题并不成立。因为劣于某个纯策略的任何策略, 必然劣于赋予此占优纯策略概率 1 的混合策略。

表 5.11

1\2	L	M	R
U	3, 1	4, 2	1, 4
D	2, 4	1, 2	3, 1

在这个博弈中,所有人的每个策略都不劣于其他纯策略。假设在参与者 2 的混合策略中, $\sigma_2(L) = \frac{1}{2}$, $\sigma_2(R) = \frac{1}{2}$ 。不管对手使用策略 U 还是策略 D ,这一混合策略的期望值都是 2.5,高于参与者 2 使用纯策略 M 时所能获得的

的收益。

下面的定理明确了混合策略 Nash 均衡和反复剔除严格劣于混合策略的策略之间的关系。

定理 5.5 设 $\Gamma^m = \langle N, \Delta, u^m \rangle$ 是有限标准式博弈,它有混合策略 Nash 均衡 σ^* 。如果在混合策略 σ_i^* 下,策略 s_i 被以正概率使用,则在反复剔除严格劣于混合策略的策略的过程中, s_i 存活下来。

我们把这一定理的证明留做习题。在计算混合策略均衡时,这一定理十分有用。有了它,我们不必对所有策略上的概率分布求最优反应,而只需对在反复剔除过程中存活下来的策略上的概率分布,求最优反应。

5.6 计算 Nash 均衡

我们虽然能够明确 Nash 均衡存在的充分条件,但是,计算博弈的均衡所需要的往往是技巧。事实上,国际象棋等博弈存在 Nash 均衡,但是没有人计算出其均衡策略。有一些技巧和算法能便利我们找到均衡。

5.6.1 有限博弈中的纯策略 Nash 均衡

我们给出一个检验过程,用它检验某一策略组合是否是均衡。对有限博弈,找到所有的纯策略 Nash 均衡的一种方法,是逐个检验每个策略组合 $s' \in S$ 是否是 Nash 均衡。我们从策略组合 $s' = (s'_1, \dots, s'_n)$ 开始,依次回答以下问题:

(1) 保持 $s'_2, \dots, s'_n$ 固定,是否存在策略 s''_1 使得 $u_1(s''_1, s'_{-1}) > u_1(s'_1, s'_{-1})$? 如果存在这样的 s''_1 ,则 s' 不是 Nash 均衡。如果不存在,则继续下面的过程。

(2) 保持 $s'_1, s'_3, \dots, s'_n$ 固定,是否存在策略 s''_2 使得 $u_2(s''_2, s'_{-2}) > u_2(s'_2, s'_{-2})$? 如果存在这样的 s''_2 ,则 s' 不是 Nash 均衡。如果不存在,则继续下面的过程。

⋮

(i) 保持 $s'_1, \dots, s'_{i-1}, s'_{i+1}, \dots, s'_n$ 固定,是否存在策略 s''_i 使得 $u_i(s''_i, s'_{-i}) > u_i(s'_i, s'_{-i})$? 如果存在这样的 s''_i ,则 s' 不是 Nash 均衡。如果不存在,则继续下面的过程。

⋮

(n) 保持 $s'_1, \dots, s'_{n-1}$ 固定,是否存在策略 s''_n 使得 $u_n(s''_n, s'_{-n}) > u_n(s'_n, s'_{-n})$? 如果存在这样的 s''_n ,则 s' 不是 Nash 均衡。如果不存在,则 s' 是 Nash 均衡。

对 S 中的每个策略组合重复上述算法。

在两人有限博弈中,即可以用矩阵表示的博弈中,这种算法很方便。从矩阵的某个单元格开始,验证在同一列中是否存在另一个单元格,能使行参与者有更高的收益。如果存在这样的单元格,则初始策略组合不是 Nash 均衡。如果不存在,则交换行和列,重复这一步骤。

在表 5.12 中,假设 (t, l) 是纯策略 Nash 均衡。给定 $s_2 = l$, 参与者 1 在 $u_1(t, l) = 5$, $u_1(m, l) = 2$ 和 $u_1(b, l) = 5$ 之间做出选择。他的最优反应为 $b_1(l) = \{t, b\}$ 。在 $s_1 = t$ 时,参与者 2 在收益 4、3 和 2 之间做出选择, $b_2(t) = \{l\}$ 。因此, (t, l) 是 Nash 均衡。由于 $b_2(t)$ 中只有一个元素,所以, (t, l) 是包含策略 t 的唯一的纯策略 Nash 均衡。由于 $b_1(l) = \{t, b\}$, 所以, (m, l) 不是 Nash 均衡,因为在参与者 2 选择 l 时,参与者 1 会偏离到 t 或 b 。假设 (m, c) 是 Nash 均衡。由于 $b_1(c) = \{m\}$ 且 $b_2(m) = \{c\}$, 所以,这一假设成立。由于 $r \notin b_2(m)$, 所以包含策略 m 的唯一的纯策略 Nash 均衡是 (m, c) 。假设 (b, l) 是 Nash 均衡。由于 $b_2(b) = \{r\}$, 所以这一假设不成立。由于 $b_1(c) = \{m\}$, 所以 (b, c) 不是 Nash 均衡。假设 (b, r) 是 Nash 均衡。两个参与者如果偏离这一策略组合,都无法提高自己的收益,所以,这一假设也成立。通过这一繁琐的过程,我们可以断定,这一博弈的纯策略 Nash 均衡集为 $\{(t, l), (m, c), (b, r)\}$ 。

表 5.12

1\2	l	c	r
t	5, 4	2, 3	6, 2
m	2, 5	3, 6	5, 5
b	5, 2	0, 3	7, 4

5.6.2 有连续策略空间的博弈中的纯策略 Nash 均衡

除非策略空间是有限集,否则上述算法不起作用。^①我们有时还能用优化方法计算均衡。如果效用函数 $u_i(s)$ 二阶可导,我们就可利用微分求得最优反应对应。由于对 s_{-i} 的最优反应是 $u_i(s_i, s_{-i})$ 在 S_i 上的最大化问题的解,所以, $s_i \in b_i(s_{-i})$ 的充分条件是 $\partial u_i(s_i, s_{-i}) / \partial s_i = 0$ 和 $\partial^2 u_i(s_i, s_{-i}) / \partial s_i^2 < 0$ 。其中前一个条件是一阶条件(FOC),后一个条件是二阶条件(SOC)。^②如果 FOC 和 SOC 成立,则方程 $\partial u_i(s_i, s_{-i}) / \partial s_i = 0$ 的解给出了最优反应对应。而且,如果 $u_i(s)$ 严格凹,则满足 FOC 和 SOC 的解是唯一的,最优反应为函数。在这种情况下,下面的方程组的每个解都是 Nash 均衡:

$$s_1^* = b_1(s_{-1}^*)$$

$$s_i^* = b_i(s_{-i}^*)$$

① 在学习本节前,你可能需要温习数学附录中关于优化的讨论。

② 如果 S_i 是多维空间, $\partial u_i(s_i, s_{-i}) / \partial s_i$ 就是个向量,其中每个坐标都是对 s_i 中的一个坐标的偏导数。0 表示零向量。二阶条件要求二阶导数矩阵是负定的。

$$s_n^* = b_n(s_{-n}^*)$$

这种方法有时会相当繁复——它要求既求解最优反应函数又求解方程组。一般情况下,直接求解下面的由一阶条件构成的方程组更方便:

$$\frac{\partial u_1(s_1, s_{-1})}{\partial s_1} = 0$$

$$\vdots$$

$$\frac{\partial u_i(s_i, s_{-i})}{\partial s_i} = 0$$

$$\vdots$$

$$\frac{\partial u_n(s_n, s_{-n})}{\partial s_n} = 0$$

要确保方程组的解是 Nash 均衡,还需要对每个参与者验证二阶条件在解上是否成立。

如果上述两个方程组都有多个解,就存在多个纯策略均衡。但是,如果没有解,并不意味着不存在纯策略 Nash 均衡。FOC 和 SOC 是充分条件而不是必要条件,在许多博弈中,最优反应并不满足这两类条件。原因有许多。给定 s_{-i} ,如果收益函数不是拟凹的,或者如果它们在 S_i 上单调递增或单调递减,参与者 i 的最优反应就在 S_i 的边界上。这种最优反应被称为“角解”,一般不满足一阶条件。如果收益函数在 S_i 上不连续,FOC 方法也不起作用;在不连续的点上,我们无法求解一阶或二阶条件。

5.7 应用:利益集团献金

两个利益集团 $N = \{1, 2\}$ 试图影响政府政策。两个集团都知道,最终政策是它们捐给政府的竞选资金的函数。集团 1 最偏好的政策是 0,集团 2 最偏好的政策是 1。政府想要的政策是 $\frac{1}{2}$,但最终政策受竞选献金影响。每个集团选择献金额 $s_i \in [0, 1]$,最终政策是 $x(s_1, s_2) = \frac{1}{2} - s_1 + s_2$ 。两个集团同时做出选择,政府拿到献金以便为下次选举购买广告。^①两个利益集团在献金数量和最终政策上的效用函数分别为:

$$u_1(s_1, s_2) = -(x(s_1, s_2))^2 - s_1$$

$$u_2(s_1, s_2) = -(1 - x(s_1, s_2))^2 - s_2$$

① 这一博弈类似经济学中的全支付(all-pay)拍卖。在本书后半部分中,我们将讨论拍卖。

把政策函数代入到两个效用函数中,得到:

$$u_1(s_1, s_2) = -\left(\frac{1}{2} - s_1 + s_2\right)^2 - s_1$$

$$u_2(s_1, s_2) = -\left(1 - \left(\frac{1}{2} - s_1 + s_2\right)\right)^2 - s_2$$

一阶条件为:

$$FOC_1: 2\left(\frac{1}{2} - s_1 + s_2\right) - 1 = 0$$

$$FOC_2: 2\left(1 - \left(\frac{1}{2} - s_1 + s_2\right)\right) - 1 = 0$$

求解 FOC_1 , 把 s_1 表示为 s_2 的函数, 得到最优反应:

$$b_1(s_2) = s_2$$

求解 FOC_2 , 得到最优反应:

$$b_2(s_1) = s_1$$

在这个博弈中, Nash 均衡并不唯一。任何一对献金额, 只要满足 $s_1 = s_2$, 就是 Nash 均衡。我们把这一博弈的纯策略 Nash 均衡集写为:

$$\{(s_1, s_2) \in [0, 1]^2: s_1 = s_2\}$$

这一结论有很直接的解释。任何一对等值献金都是 Nash 均衡, 并且在所有这些均衡中, 政策结果都是 $\frac{1}{2}$ 。没有哪个利益集团单方面偏离这一均衡, 因为增加一元钱献金所实现的边际收益(向自己想要的方向拉动政策)正好与减少一元钱的资源的边际成本相抵消。和囚徒两难博弈一样, 这一博弈中的 Nash 均衡是无效率的。就是说, 献金者完全可以相互承诺不给政府任何捐款。但是, 这种承诺行不通, 所以在均衡中, 献金的数量是正的。

5.8 应用: 国际外部性

两个国家必须决定用于治理污染的投资规模。他们的策略分别是投资水平 $s_1 > 0$ 和 $s_2 > 0$ 。每个国家还要发生相应的成本 $c(s_i) = k_i s_i$ 。设 $k_1 < k_2$, 国家 1 能够以更低成本减少给定数量的污染。总污染水平影响两个国家的居民, 所以, 从污染治理中得到的效用决定于总投资水平, 设效用函数为 $u(s_1, s_2) = \sqrt{s_1 + s_2}$ 。每个国家的收益决定于从污染治理中得到的效用和付出的成本, 因此收益函数为:

$$\sqrt{s_1 + s_2} - k_i s_i$$

利用收益函数, 得到一阶条件:

$$\frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}} - k_1 = 0$$

$$\frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}} - k_2 = 0$$

两个国家面对的二阶条件也得到满足： $-\frac{1}{4}(s_1 + s_2)^{-\frac{3}{2}} < 0$ 。

由于 $k_1 < k_2$ ，一阶条件构成的方程组无解。假设 $s_1 > (2k_2)^{-2}$ 。则国家 2 的一阶条件中等号左边的部分始终是负的，这就意味着国家 2 的收益在其投资水平上始终递减。在图 5.4 中，国家 2 的收益被表示为国家 1 的投资水平的函数，图中画出了在 $s_1''' > (2k_2)^{-2}$ 时国家 2 的收益的递减特征。由于投资是非负的，国家 2 对国家 1 的这一投资水平的最优反应是 0。同样地，如果 $s_2 > (2k_1)^{-2}$ ，国家 1 的最优反应也是投资为 0。这些“角解”构成了各个国家的最优反应函数，所以，我们无法用一阶条件求解。

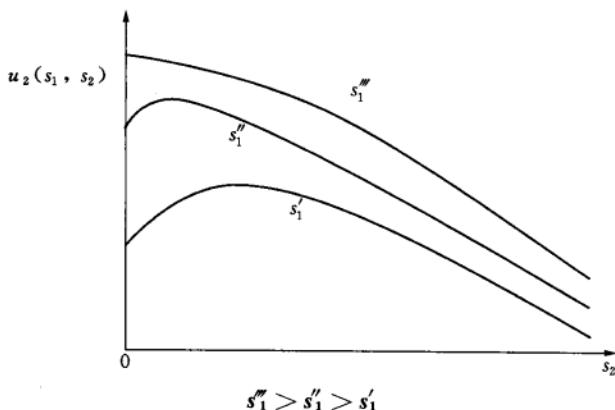


图 5.4 国家 2 的收益是国家 1 的投资水平的函数

但是，我们能用图形求出这一解。在图 5.5 中，我们画出了两个国家的最优反应函数。纵轴为国家 2 的策略和它对国家 1 的最优反应；横轴为国家 1 的策略和它的最优反应。虚线为国家 2 的最优反应函数，它在 s_1 上递减直至 $s_1 > (2k_2)^{-2}$ ，此时它为 0。实线是国家 1 的最优反应函数，它亦递减且在 $s_2 = (2k_1)^{-2}$ 时为 0。从图 5.5 上可以清楚地看出，这两个最优反应曲线的唯一交点是 $s_1 = (2k_1)^{-2}$ 和 $s_2 = 0$ 。

如果 $k_2 < k_1$ ，国家 2 就是低成本国家。此时，Nash 均衡策略为 $s_1 = 0$ 和 $s_2 = (2k_2)^{-2}$ 。这一博弈的唯一的 Nash 均衡是高成本国家完全免费搭乘低成本国家便车。低成本国家知道高成本国家不会投资，在这一认识基础上，它选择自己的最优投资水平。

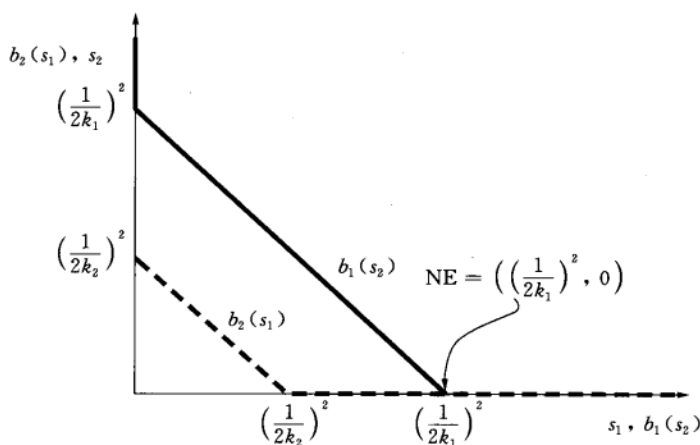


图 5.5 外部性博弈中的最优反应函数

5.9 利用约束条件下的优化方法, 计算均衡

当博弈中的最优反应函数映射到策略空间的边界时, 我们可以用约束条件下的最大化技术, 如 Kuhn-Tucker 优化方法, 求 Nash 均衡的必要条件。^①为说明这种方法, 我们再来分析外部性博弈。现在加入约束条件 $s_1 \geq 0$ 和 $s_2 \geq 0$ 。我们把这两个条件通过拉格朗日乘数 λ_1 和 λ_2 , 引入到优化问题中。参与者 i 最大化:

$$\sqrt{s_1 + s_2} - k_i s_i + \lambda_i s_i$$

Nash 均衡的必要条件是一阶条件:

$$\frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}} - k_1 + \lambda_1 = 0$$

$$\frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}} - k_2 + \lambda_2 = 0$$

和“松弛性”条件:

$$\lambda_1 s_1 = 0$$

$$\lambda_2 s_2 = 0$$

及约束条件:

$$\text{对所有的 } i, \lambda_i \geq 0$$

$$\text{对所有的 } i, s_i \geq 0$$

^① 你可能需要重温数学附录中关于约束条件下的优化的介绍。

Nash 均衡是满足约束条件的四个方程的解。根据前两个方程,
 $\lambda_i = k_i - \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}}$ 。把这一结果代入到松弛性条件中,得到:

$$s_1 \left[k_1 - \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}} \right] = 0$$

$$s_2 \left[k_2 - \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}} \right] = 0$$

把这一结果代入到关于 λ 的约束条件中,得到:

$$k_1 \geq \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}}$$

$$k_2 \geq \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}}$$

和前面的结论一样,我们没有得到两个国家投资数量都大于零的解。这种解要求一阶条件中括号内的项为 0,而这是不可能的。策略组合 $s_1 = s_2 = 0$ 也不是均衡,因为条件 $k_i \geq \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}}$ 没有被满足。假设 $s_1 = 0$ 且 $s_2 > 0$ 。根据第二个一阶条件, $k_2 = \frac{1}{2}(s_1 + s_2)^{-\frac{1}{2}}$,这就违背了 λ_1 上的非负性约束。所以,唯一可能的 Nash 均衡是 $s_1 > 0$ 且 $s_2 = 0$ 。根据第一个一阶条件, $s_1^* = (2k_1)^{-2}$ 。这正是上一节中的结论。

5.10 证明 Nash 均衡的存在性 **

我们给出 Nash 定理的严格证明。首先是有关数学概念,其中最重要的概念是不动点。直观地说,一个函数(或一个对应)的不动点,是定义域中的一个点,它的像是它自身。正式地说,给定对应 $c: A \rightarrow A$, 不动点 $x^* \in A$ 是满足 $x^* \in c(x^*)$ 的点。如果 $c(\cdot)$ 是函数,则不动点 x^* 是满足 $x^* = c(x^*)$ 的点。在 Nash 均衡策略组合 s^* 上,对每个 $i \in N$, $s_i^* \in b_i(s_{-i}^*)$, 所以,它是最优反应对应:

$$b(s) = (b_1(s_{-1}), \dots, b_i(s_{-i}), \dots, b_n(s_{-n}))$$

的不动点。证明 Nash 均衡的存在性,是要确定最优反应对应是否存在不动点。幸运的是,有一些数学概念可用于确定不动点存在性的充分条件。这一充分条件,被称为不动点定理。如果博弈的最优反应对应的性质满足不动点定理的前提条件,则 Nash 均衡存在。

不动点的存在性的证明,要用到许多概念。第一个概念是凸值对应。

定义 5.11 对应 $c: A \rightarrow A$ 是凸值的,如果对每个 $a \in A$, $c(a)$ 是 A 的凸子集。

对每个 $x \in X$, 如果 $y, z \in c(x)$ 且对任意的 $\lambda \in [0, 1]$, $\lambda y + (1 - \lambda)z \in c(x)$, 则对应 $c(\cdot)$ 是凸值的。

另一个有用的定义是集合的上反像(upper inverse)。给定集合 B , 其上反像是定义域中被对应 $c(\cdot)$ 映射到 B 的子集的点集。

定义 5.12 给定对应 $c: A \rightarrow A$, 集合 $B \subseteq A$ 的上反像是集合 $c^+(B) = \{x \in A: c(x) \subset B\}$ 。

上反像对定义对应的一个连续性概念——上半连续性——起着重要作用。

定义 5.13 对应 $c: A \rightarrow A$ 是上半连续的, 如果对每个开集 $O \subseteq A$, 上反像 $c^+(O)$ 是开集。

定义 5.13 常常很难运用, 而闭图的存在性相对容易确定。

定义 5.14 对应 $c: A \rightarrow A$ 有闭图, 如果对任意两个点列 $x^n \rightarrow x \in A$ 和 $y^n \rightarrow y \in A$ 且对每个 n , 有 $x^n \in A$ 和 $y^n \in c(x^n)$, 则有 $y \in c(x)$ 。

下面的定理(其证明留做习题)指出, 有闭图的对应是上半连续的。

定理 5.6 如果 A 是紧集, 且对应 $c: A \rightarrow A$ 有闭图, 则这一对应是上半连续的。

闭图条件背后的原理并不难。设在 A 中有两个点列 x^n 和 y^n , 使得对每个 n , $y^n \in c(x^n)$ 并且这两个点列分别收敛于 A 中的点 x 和 y 。如果 c 有闭图, 则 $y \in c(x)$ 。并且, 对有闭图的对应, 集合 $\{(x, y) \in A^2: y \in c(x)\}$ 在空间 A^2 中是闭的。对这些概念的更全面的讨论, 见 Border(1989)。

我们说过, 对每个 $a \in A$, 如果对应 $c: A \rightarrow A$ 是单值的, 它就是个函数。如果单值对应是上半连续的, 则它是连续函数。

要证明定理 5.3, 我们需要用到 Brouwer 不动点定理。

定理 5.7 (Brouwer) 设 $A \subset \mathbb{R}^d$ 是紧的凸集。如果 $f: A \rightarrow A$ 是连续函数, 则 f 在 A 中有不动点。

要证明定理 5.4, 我们需要用到角谷(Kakutani)不动点定理。

定理 5.8 (Kakutani) 设 $A \subset \mathbb{R}^d$ 是紧的凸集。如果对应 $c: A \rightarrow A$ 满足以下条件:

- (1) 对每个 $x \in A$, $c(x)$ 是非空集;
 - (2) $c(\cdot)$ 是凸值的;
 - (3) $c(\cdot)$ 是上半连续的;
- 则 $c(\cdot)$ 在 A 中有不动点。

Border(1989)中给出了这些定理的各种证明。要证明 Nash 均衡的存在性, 如果使用定理 5.7, 我们只需证明 $b(s)$ 是个连续函数; 如果使用定理 5.8, 我们只需证明 $b(s)$ 是满足条件(1)–(3)的对应。Berger(1997)给出的一个定理相当有用, 能证明 $b(s)$ 是非空的和上半连续的。

定理 5.9 (Berger 的最大值定理) 设 $X \subset \mathbb{R}^d$ 和 $M \subset \mathbb{R}^z$ 是紧的凸集。设函数 $f(x, m): X \times M \rightarrow \mathbb{R}$ 在 x 和 m 上连续。定义对应 $c: M \rightarrow X$ 为:

$$c(m) = \arg \max_{x \in X} \{f(x, m)\}$$

则对每个 $m \in M$, $c(m)$ 是非空的和上半连续的。

对定理 5.9 中 $c(\cdot)$ 上半连续的结论, 我们可以这样理解: 设向量 m 是上述优化问题中的参数向量。如果参数向量列 m^n 收敛于 m , 则对收敛于 x 的任意最优政策列 $x^n \in c(m^n)$, $x \in c(m)$ 。

我们证明定理 5.3。

定理 5.3 的证明: 对每个 $i \in N$, 设 S_i 是 $\mathbb{R}^d$ 的凸子集; 对每个 $i \in N$, 对每个 $s_{-i} \in S_{-i}$, 设 $u_i(s): S \rightarrow \mathbb{R}$ 在 s_i 上连续和严格拟凹。定义对应 $b_i(s_{-i}): S_{-i} \rightarrow S$ 为:

$$b_i(s_{-i}) = \arg \max_{s_i \in S_i} \{u_i(s_i, s_{-i})\}$$

根据定理 5.9, 对每个 $s_{-i} \in S_{-i}$ 它是非空的, 且是上半连续的。严格拟凹的效用函数可以表示为严格凸的偏好序, 所以根据定理 2.4, 对每个 $s_{-i} \in S_{-i}$, $b_i(s_{-i})$ 是单元素集。而单值上半连续的对应是连续函数的事实(见习题 5.15), 意味着对每个 $i \in N$, $b_i(s_{-i})$ 是从 S_{-i} 到 S_i 的连续函数。定义 $b(s_1, \dots, s_n) = (b_1(s_{-1}), \dots, b_n(s_{-n}))$, 我们得到函数:

$$b(s): S \rightarrow S$$

向量 s_{-i} 是 s 的连续函数(投影), 我们可以构造出复合函数 $\tilde{b}_i(s) = b_i(s_{-i}(s))$ 。函数 $b(s)$ 因此是连续的, 因为: (1) $b_i(\cdot)$ 连续, (2) 连续函数的复合 $\tilde{b}_i(s)$ 连续, (3) 连续函数的积 $(\tilde{b}_1(s), \dots, \tilde{b}_n(s))$ 连续。根据 Brouwer 不动点定理, 这一映射有不动点, 即存在 s^* , $s^* = b(s^*)$ 。这就是说, 对每个 $i \in N$, $b_i(s_{-i}^*) = s_i^*$, 所以 s^* 是 Nash 均衡。□

定理 5.4 的证明与此类似。它告诉我们, 在有限博弈的混合形式中, 最优反应对应满足角谷不动点定理中的前提条件。

定理 5.4 的证明: 在标准式博弈 Γ 中, 对每个 $i \in N$, 设策略集 S_i 是有限集。于是, 在混合博弈 Γ^m 中, Δ_i 是有限维欧几里得空间的紧的凸子集。根据定义 5.9, $u_i(\sigma_i, \sigma_{-i})$ 是线性的并因此在 σ 上连续。定义 $b_i: \Delta_{-i} \rightarrow \Delta_i$ 为:

$$b_i(\sigma_{-i}) = \arg \max_{\sigma_i \in \Delta_i} \{u_i(\sigma_i, \sigma_{-i})\}$$

根据定理 5.9, 对每个 $\sigma_{-i} \in \Delta_{-i}$, $b_i(\sigma_{-i})$ 是非空的, 且上半连续。由于对任何的 σ_{-i} , $u_i(\sigma_i, \sigma_{-i})$ 是线性的, 所以给定任意的 $\lambda \in (0, 1)$, 如果 $u_i(\sigma'_i, \sigma_{-i}) = u_i(\sigma''_i, \sigma_{-i})$, 则 $u_i(\lambda \sigma'_i + (1-\lambda)\sigma''_i, \sigma_{-i}) = u_i(\sigma'_i, \sigma_{-i})$ 。这就是说, $b_i(\sigma_{-i})$ 是凸值的。定义对应:

$$b(\sigma): \Delta \rightarrow \Delta$$

为 $b(\sigma_1, \dots, \sigma_n) = (b_1(\sigma_{-1}), \dots, b_n(\sigma_{-n}))$ 。根据上述结论, 这一对应满足角谷不动点定理的前提条件。所以, 存在混合策略组合 σ^* , 满足条件 $\sigma^* \in b(\sigma^*)$ 。这样

的策略组合是 Γ^m 的 Nash 均衡, 是 Γ 的混合策略 Nash 均衡。□

5.11 比较静态

博弈论模型具有的一个特点, 使得它们特别适用于应用研究, 这就是: 对政治问题中的外生性质如何影响其内生性质, 博弈论模型能够给出解释。这种关系被称为比较静态。^①我们看一个简单的竞选模型。在竞选中, 候选人为获得公职, 同时选择竞选广告上的支出数量。假设有两个候选人, 称为 1 和 2。候选人 $i \in \{1, 2\}$ 选择一个非负实值支出水平 $s_i \in \mathbb{R}_+$ 。无论哪个人最终当选, 每个人都要承担由支出 s_i 导致的成本。我们用 $c_i(s_i)$ 表示支出水平 s_i 给候选人 i 造成的成本。在竞选中胜出所产生的效用为 1, 落选所产生的效用为 0。 $p(s_1, s_2)$ 表示在支出水平分别为 s_1 和 s_2 时, 候选人 1 胜出的概率。两位候选人的收益因此为:

$$u_1(s_1, s_2) = p(s_1, s_2) - c_1(s_1)$$

$$u_2(s_1, s_2) = 1 - p(s_1, s_2) - c_2(s_2)$$

首先设 $c_i(s_i) = ks_i^2$ 且 $p(s_1, s_2) = \max \left\{ 0, \min \left\{ 1, \frac{1}{2} + \beta(s_1 - s_2) \right\} \right\}$, 其中 k 和 β 是严格为正的参数。为保证有内解, 我们假设 $\beta^2 < 2k$ 。设 s_1^* 和 s_2^* 表示均衡解, 它满足条件 $p(s_1^*, s_2^*) \in (0, 1)$ 。微分得到一阶条件:

$$\beta = 2ks_1$$

$$\beta = 2ks_2$$

均衡解为:

$$s_1^* = \frac{\beta}{2k}$$

$$s_2^* = \frac{\beta}{2k}$$

为突出参数和均衡行为间的关系, 我们将均衡解写为 $s_i^*(\beta, k) = \beta/2k$ 。为说明均衡解如何随外生参数变化, 我们将其微分, 得到:

$$\frac{\partial s_i^*(\beta, k)}{\partial \beta} = \frac{1}{2k}$$

$$\frac{\partial s_i^*(\beta, k)}{\partial k} = -\frac{\beta}{2k^2}$$

① 不熟悉微积分的读者, 在学习本节前, 应该重温数学附录中的微积分部分。

就是说,均衡广告水平在广告的边际产出上递增,在广告的边际成本上递减。在这个例子中,一阶条件所产生的最优反应函数并不决定于对手的选择变量;一般情况却并非如此。正是因为这个原因,在这个例子中,比较静态分析十分简单。而在一般情况中,一阶条件以隐函数形式给出了每个人的最优反应对应,均衡则是最优反应对应的不动点。原则上,我们并不愿意对 $p(\cdot, \cdot)$ 和 $c(\cdot)$ 的函数形式施加很强的假设。那么,如果没有具体的函数形式,我们如何了解均衡点的比较静态特征呢? 如果我们所选择的函数形式无法产生关于 s_i^* 的闭式解,即无法将 s_i^* 表示为参数的显函数形式,又如何进行分析呢? 隐函数定理为这类问题提供了答案。在广告博弈中,可以用 $p(s_1, s_2; \beta)$ 把候选人 1 的胜出概率表示为选择变量和参数的函数;用 $c_1(s_1; k_1)$ 和 $c_2(s_2; k_2)$ 把成本表示为选择变量和参数 k_1 或 k_2 的函数。我们习惯上用分号把内生变量与外生变量区分开来。上面的例子是这类模型的特殊情况。现在的一阶条件是:

$$\begin{aligned} \frac{\partial p(s_1, s_2; \beta)}{\partial s_1} - \frac{\partial c_1(s_1; k_1)}{\partial s_1} &= 0 \\ \frac{\partial p(s_1, s_2; \beta)}{\partial s_2} - \frac{\partial c_2(s_2; k_2)}{\partial s_2} &= 0 \end{aligned} \quad (5.1)$$

为避免一些技术上的细节问题,我们假定方程组(5.1)产生单值最优反应;就是说,我们假定存在最优反应函数。^①在探讨均衡前,我们分析最优反应如何随外生参数和对手的策略而变化。隐函数定理指出,只要某些条件得到满足,就能用微分技术分析在 s_2 、 β 或 k_1 的值发生微小变化时,候选人 1 的最优选择是如何变化的。数学附录中给出了这一定理的一般表述。在现在所探讨的问题中,我们做如下处理。 $b_1(\cdot, \cdot, \cdot, \cdot)$ 是参与者 1 的最优反应,它把 $(s_2; \beta, k_1, k_2)$ 映射到 s_1 。设在参数向量 $(s_2; \beta, k_1, k_2)$ 上,函数 $p(\cdot, \cdot; \cdot)$ 、 $c_1(\cdot; \cdot)$ 和 $c_2(\cdot; \cdot)$ 连续可微,并且在 $s_1 = b_1(s_2; \beta, k_1)$ 时, $\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} \neq \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}$ 。根据隐函数定理,在向量 $(s_2; \beta, k_1)$ 的某邻域中,函数 b_1 可微且有如下二阶偏导数:

$$\begin{aligned} \frac{\partial s_1}{\partial s_2} &= - \frac{\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_2}}{\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}} \\ \frac{\partial s_1}{\partial \beta} &= - \frac{\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial \beta}}{\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}} \end{aligned}$$

① 你可以找出使这一假设成立,原始函数需满足的充分条件。参考与定理 5.3 有关的一些结论。

$$\frac{\partial s_1}{\partial k_1} = - \frac{\frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial k_1}}{\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}}$$

在 $s_1 = b_1(s_2; \beta, k_1)$ 时, $\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} \neq \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}$ 的条件,保证了在上述各个式子中,分母不为 0。隐函数定理能使我们分析参数如何影响一阶条件。作为习题,请你对 $b_2(\cdot; \cdot, \cdot)$ 进行有关的比较静态分析。即使放松一阶条件产生唯一的最优反应的假定,比较静态分析也不是难事。只要 $\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} \neq \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}$, 一阶条件的解在局部仍是唯一的;在一阶条件产生的任何解的邻域中,不存在其他的解。因此,比较静态分析能反映在最优反应对应的每个变量上的局部变化。

但是,有一点必须注意。上述分析探讨的是最优反应如何随参数而变化,而没有告诉我们均衡行为是如何变化的。我们现在探讨这后一问题。

5.11.1 均衡效应*

均衡策略组合 (s_1^*, s_2^*) 是方程组 (5.1) 的解。^①为探讨参数 (β, k_1, k_2) 的变化如何影响均衡,我们对这个方程组应用隐函数定理。在上一节中,我们只是对各个最优反应应用隐函数定理。要分析均衡的比较静态特征,需要用到矩阵代数中的工具。雅可比矩阵为:

$$\begin{vmatrix} \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_2} \\ \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_2} - \frac{\partial^2 c_2(s_2; k_2)}{\partial s_2 \partial s_2} \end{vmatrix}$$

在前面,我们要求 $\frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} \neq \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1}$, 这里则要求上述矩阵的行列式是非奇异的(即行列式不等于零)。^②设函数 $p(\cdot, \cdot; \cdot)$ 、 $c_1(\cdot; \cdot)$ 和 $c_2(\cdot; \cdot)$ 连续可微,且在参数向量 (β, k_1, k_2) 上,向量 (s_1^*, s_2^*) 是方程组 5.1 的解。设在这一解上,雅可比行列式是非奇异的。隐函数定理指出,存在向量 (β, k_1, k_2) 的邻域 B

① 本节用到线性代数和多变量微分中的概念。

② 这是个 2×2 矩阵。给定 2×2 矩阵 $\begin{bmatrix} a & b \\ c & d \end{bmatrix}$, 行列式为 $\det = ad - cb$, 这一行列式的倒数为

$\frac{1}{\det \begin{bmatrix} d & -a \\ -b & a \end{bmatrix}}$ 。如果矩阵更大,计算会更繁复。

和解 (s_1^*, s_2^*) 的邻域 A ,在这些邻域上,方程组 5.1 的解由在参数上可微的函数 $s: B \rightarrow A$ 给出,这一函数的一阶导数为:

$$\begin{aligned} \begin{bmatrix} \frac{\partial s_1}{\partial \beta} \\ \frac{\partial s_2}{\partial \beta} \end{bmatrix} &= - \begin{bmatrix} \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_2} \\ \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_2} - \frac{\partial^2 c_2(s_2; k_2)}{\partial s_2 \partial s_2} \end{bmatrix}^{-1} \begin{bmatrix} \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial \beta} \\ \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial \beta} \end{bmatrix} \\ \begin{bmatrix} \frac{\partial s_1}{\partial k_1} \\ \frac{\partial s_2}{\partial k_1} \end{bmatrix} &= - \begin{bmatrix} \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_2} \\ \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_2} - \frac{\partial^2 c_2(s_2; k_2)}{\partial s_2 \partial s_2} \end{bmatrix}^{-1} \begin{bmatrix} \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1} \\ 0 \end{bmatrix} \\ \begin{bmatrix} \frac{\partial s_1}{\partial k_2} \\ \frac{\partial s_2}{\partial k_2} \end{bmatrix} &= - \begin{bmatrix} \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_1} - \frac{\partial^2 c_1(s_1; k_1)}{\partial s_1 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_1 \partial s_2} \\ \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_1} & \frac{\partial^2 p(s_1, s_2; \beta)}{\partial s_2 \partial s_2} - \frac{\partial^2 c_2(s_2; k_2)}{\partial s_2 \partial s_2} \end{bmatrix}^{-1} \begin{bmatrix} 0 \\ \frac{\partial^2 c_2(s_2; k_2)}{\partial s_2 \partial s_2} \end{bmatrix} \end{aligned}$$

有几点需要指出。第一, $\frac{\partial s_1}{\partial k_2}$ 可能不为零的事实说明了均衡效应的重要性。前面说过, k_2 并不影响候选人 1 的最优反应。 k_2 的变化能改变最优反应 $b_2(\cdot; \cdot)$; 而最优反应上的变化会改变 s_1 的均衡解。第二, 确定比较静态变化的方向, 可能是个繁琐的过程, 特别是当所探讨的问题涉及高维策略空间时。第三, 一阶条件有多个解的可能性也会给分析增加难度。就最优反应对应而言, 比较静态变化的方向可能并不确定。但是, 如果雅可比行列式为非奇异的, 则每个均衡都有一个简单性质。这些均衡在局部是唯一的: 在每个均衡点的某个邻域中, 没有其他均衡。而且, 对参数上的微小变化, 任何新的均衡在局部也是唯一的。在存在多个均衡时, 常遇到的一个问题是, 对给定的参数向量, $\frac{\partial s_i}{\partial \beta}$ 在两个不同的均衡点上, 可能有相反的符号。下面两节探讨在哪些类型的博弈中, 这一问题能得到解决。

5.11.2 策略互补性

假设候选人 1 和候选人 2 通过选择支出水平 $s_i > 0$ 竞选公职。为简化分析, 设竞选支出的机会成本是 $c_i(s_i) = c_i s_i$, 其中, $c_i > 0$ 。候选人 1 胜出的概率 $p(s_1, s_2)$ 在 s_1 上递增, 在 s_2 上递减。在建立竞选模型时, 我们要解决一个重要问题: 在这样的竞选中, 候选人 1 的最优反应随候选人 2 的支出水平递增还是递减? 常理是, 面对着更高的 s_2 , 候选人 1 会选择更高的 s_1 ; 反之亦然。如果真是这样, 我们就说, 这两个选择变量是“策略互补”的。具有策略互补性的博弈被称为超模博弈。这种博弈尤其方便进行均衡和比较静态分析。本节用一个例子阐述这种分析背后的原理, 下一节探讨其中的技术细节。

Kahn 和 Kenney(1990)在对参议院选举的研究中, 认为竞争性是决定竞选活

动所产生的媒体覆盖率的重要因素。竞争性决定选民对竞选活动的反应方式,而竞选活动则影响竞选的竞争性。没有竞争性选举,媒体和选民就不会关注,竞选广告的边际价值就很低。另一方面,在竞争性很高时,竞选广告就有很大的影响。在 Kahn 和 Kenney 的理论中,媒体被描述为一种机制,它使得竞争性是竞选活动的函数。根据这种认识,更高的竞选支出水平可能导致更大的竞争性。在激烈竞争的竞选活动中,广告和其他竞选支出产生更大影响。这种回馈效应表明,候选人的广告支出水平具有策略互补性。因此,两个候选人的最优反应都是对手策略的严格增函数。在图 5.6 中,这种形式的博弈中的最优反应产生唯一的均衡点 (s_1^*, s_2^*) 。

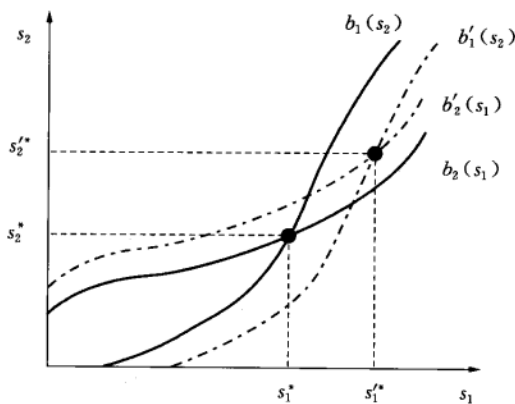


图 5.6 策略互补性

我们也能画出有多个均衡点的图。在多均衡点的图中,各个均衡点完全按顺序排列,就是说,如果 (s_1^*, s_2^*) 和 (s_1^{**}, s_2^{**}) 是两个均衡点且 $s_1^* > s_1^{**}$, 则 $s_2^* > s_2^{**}$ (如果 $s_2^* > s_2^{**}$, 则 $s_1^* > s_1^{**}$)。要理解这一点,你不妨在最优反应严格递增的前提下,看能否画出相反的情况。

我们还可以更深入地讨论 Kahn 和 Kenney 关于竞争性的探讨。设在竞选开始时,存在某个外生的竞争程度。影响竞争程度的因素可能包括职位的重要性、媒体环境和选民的重视程度。设有两种选举环境,其中一个的竞争程度更高。给定 s_i , 我们能够想象得到,在更激烈的竞选中, $b_{-i}(s_i)$ 更高。在图 5.6 中, $b'_i(s_{-i})$ 为竞争程度更高的选举中的最优反应,而 $b_i(s_{-i})$ 表示竞争程度低的选举中的最优反应。从图 5.6 中可以看出,竞争激烈的选举活动的均衡竞选支出 $(s_1'; s_2')$, 高于竞争程度低的竞选中的均衡支出 (s_1^*, s_2^*) 。

这一比较静态结果产生自策略互补性假定。只要两个最优反应是递增的,则两条最优反应曲线同时向上移动,导致更高的交点。但是,如果其中一位候选人的最优反应向下倾斜,竞争程度对均衡的影响并不确定。只要候选人的支出水平在策略上互补,均衡结果就不依赖关于函数形式或参与者偏好的任何假设。如果要使用隐函数定理,需要最优反应函数连续和可微;在这里,我们没有这种需要。

当然了,这一切的关键是找出模型中与互补性相一致的前提条件。在下面的技术性部分中,我们介绍超模博弈的基本理论,给出具有单调比较静态特征的均衡的存在性所需的充分条件。

5.11.3 超模和单调比较静态*

Nash 定理要求最优反应对应连续且策略集是紧集。在这些前提下,角谷不动点定理证明了均衡的存在性。问题是,这些前提常常要求我们对博弈模型施加相当强的假设。但是,如果最优反应满足前一节中所讨论的单调性条件,就可以用另一种不动点定理证明 Nash 均衡的存在性。而且,在最优反应满足这些条件时,均衡集的比较静态分析可谓轻而易举。和隐函数定理进行的分析相比,这种均衡分析更简单更直接。在本节中,我们概括适用于超模博弈的一些定理。有了这些定理,我们就能证明在最优反应对应不连续时均衡的存在性,并能直接得到比较静态结论。

本节中的概念适用于任意偏序。但为使内容具体,我们只探讨 $\mathbb{R}^n$ 子集上的自然偏序,即“ $\geq$ ”。首先给出一些重要定义。给定两个实数 $x, y \in \mathbb{R}$, 它们的并(join)为 $x \vee y = \max\{x, y\}$, 它们的交(meet)为 $x \wedge y = \min\{x, y\}$ 。这些定义可以扩展到向量上。对 $\mathbb{R}^n$ 中任意两个向量 x 和 y , 它们的并为 $x \vee y = \{\max\{x_1, y_1\}, \dots, \max\{x_n, y_n\}\}$, 交为 $x \wedge y = \{\min\{x_1, y_1\}, \dots, \min\{x_n, y_n\}\}$ 。一个集合,如果包含其所有元素的并和交,则为格。

定义 5.15 集合 A 是格,如果对每个 $x, y \in A$, $x \vee y \in A$ 且 $x \wedge y \in A$ 。

区间是格,区间积是格,但是类似 $\{x \in \mathbb{R}^3: x_1 + x_2 + x_3 \leq 1\}$ 的集合,即单纯形,不是格。^①形象地说,正方形(区间的积)是格,三角形不是格。^②

本节探讨的是由选择变量集 X 和外生参数集 P 所定义的单人决策或多人决策的理论问题。设这两个集合是格。决策者的目标函数值决定于这两类变量:
 $f: X \times P \rightarrow \mathbb{R}$ 。

定义 5.16 给定函数 $f: X \times P \rightarrow \mathbb{R}$, 如果对所有的 $z, z' \in X \times P$,

$$f(z) + f(z') \leq f(z \vee z') + f(z \wedge z')$$

则此函数在 (x, p) 上是超模。

验证这一条件有时有些难度。而另一个更直观的概念,递增差,常常更容易验证。

定义 5.17 给定函数 $f: X \times P \rightarrow \mathbb{R}$, 如果对所有的 $p \leq p'$ 和 $x \leq x'$,

$$f(x', p') - f(x, p') \geq f(x', p) - f(x, p)$$

① Topkis(1998)讨论了如何把非格转化为格。

② 这里不探讨离散集合或者其他非凸集。但是,格的定义和下面的一些结论适用于非凸集。

则此函数在 (x, p) 上有递增差。

对于有递增差的函数, x 的变化产生的边际效应在 p 上递增。换句话说, 递增差的概念概括了互补性思想。和超模相比, 递增差常常更容易在模型的具体背景下予以诠释。因此, 为分析方便, 我们需要建立递增差和超模两个概念之间的等价性。

定理 5.10 (Topkis) 函数 $f: X \times P \rightarrow \mathbb{R}$ 在 (x, p) 上有递增差, 当且仅当它在 (x, p) 上超模。

函数 $f: X \times P \rightarrow \mathbb{R}$ 如果二阶可导, 则对 X 或 P 或 X 和 P 的任意坐标 z_1 和 z_2 , 超模性(和递增差)等价于 $\frac{\partial^2 f}{\partial z_1 \partial z_2} \geq 0$ 。

超模性在几种数学运算下都能得以保持, 这就使得超模模型容易处理。

定理 5.11 设 X 是格。以下命题都成立:

- (1) $f(x)$ 如果在 X 上超模且 $\alpha > 0$, 则 $\alpha f(x)$ 在 X 上超模。
- (2) $f(x)$ 和 $g(x)$ 如果在 X 上超模, 则 $f(x) + g(x)$ 在 X 上超模。
- (3) 如果 $f_t(x)$ 是 X 上的超模函数列, 其中 $t = 1, 2, \dots$ 并且对每个 $x \in X$, $f(x) = \lim_{k \rightarrow \infty} f_k(x)$, 则 $f(x)$ 在 X 上超模。
- (4) 如果 $F(\omega)$ 是集合 Ω 上的分布函数且对每个 $\omega \in \Omega$, $g(x, \omega)$ 在 X 上超模, 则 $f(x) = \int_{\Omega} g(x, \omega) dF(\omega)$ 在 X 上超模。

现在探讨超模对个人优化问题的意义。其中最重要的意义是, 超模意味着最优解有单调比较静态特征, 即 $\arg \max_{x \in X} f(x, p)$ 的每个元素都是 p 的增函数。

定理 5.12 (Topkis) 如果 $f: X \times P \rightarrow \mathbb{R}$ 在 (x, p) 超模且 $g(p) = \arg \max_{x \in X} f(x, p)$, 则对参数集的每个子集 $B \subset P$, 如果 $g(p)$ 非空, 则 $g(p)$ 的每个坐标在 p 上递增。

这一定理没有明确最优解是否存在。定理 2.3 指出, 对紧集上的下连续的序, 最优集非空。偏好的下连续性对应着目标函数的上半连续性。

定义 5.18 函数 $f: X \rightarrow \mathbb{R}$ 在点 $x_0 \in X$ 上半连续, 如果对所有的 $\varepsilon > f(x_0)$, 存在 $\delta > 0$ 使得对所有的 $x \in B(x_0, \delta)$ (以 x_0 为圆心的半径为 δ 的开球), $\varepsilon \geq f(x)$ 。函数在 X 上半连续, 如果它在每个点 $x_0 \in X$ 上半连续。

函数在点 x_0 的上半连续性, 要求 $\liminf_{x \rightarrow x_0} f(x) \leq f(x_0)$ 。利用习题 5.21 中的结论、定理 2.3 以及 Topkis 的结论, 我们得到以下定理。

定理 5.13 $f: X \times P \rightarrow \mathbb{R}$ 如果在 (x, p) 上超模, 且对每个 p , 它在 x 上上半连续, 则对每个 $p \in P$, $g(p) = \arg \max_{x \in X} f(x, p)$ 非空, 并且 $g(p)$ 的每个元素在 p 上递增。

在定理 5.4 的证明中, 我们强调了为使最优反应具有不动点定理所需的性质, 而要求目标函数具有的特征。Tarsky 不动点定理使我们能利用最优反应的单调性, 证明 Nash 均衡的存在。根据定理 5.13, 为超模和上半连续的效用函数的博弈至少有一个 Nash 均衡。我们现在给出 Tarsky 定理。

定理 5.14 (Tarsky) $f(x)$ 如果是从紧格 X 到 X 自身的递增函数, 则存在至少一个不动点 x^* 使得 $f(x^*) = x^*$ 。

图 5.7 在 X 为单维空间的情况中, 展示了 Tarsky 不动点定理。只要函数 $f(x)$ 是递增的, 它的不连续性并不意味着 $f(x)$ 和 45° 线不相交。角谷定理把 Brouwer 定理推广至对应情况中, Zhou(1994) 则把 Tarsky 定理推广至对应情况中。这里并不介绍 Zhou 的结论, 因为这需要引入更多符号。我们直接给出针对超模博弈的结论。

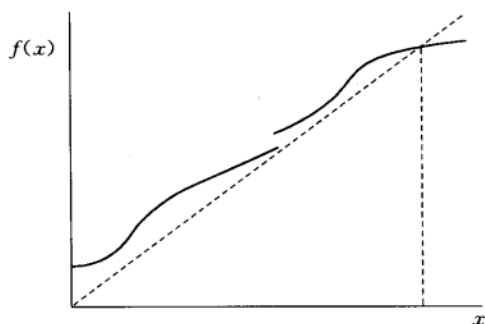


图 5.7 Tarsky 不动点定理

定义 5.19 一般的超模博弈是标准式博弈 $\langle N, \{S_i, u_i(\cdot, \dots, \cdot)\}_{i \in N} \rangle$, 如果对每个 $i \in N$, S_i 是紧格, 且对每个 $i \in N$, $u_i(\cdot, \dots, \cdot)$ 是超模和上半连续的。

我们现在给出关于一般超模博弈的主要结论。

定理 5.15 一般的超模博弈至少有一个纯策略 Nash 均衡, 并且 Nash 均衡集是个格, 包含最小的均衡策略组合和最大的均衡策略组合。

我们为什么要强调最小均衡和最大均衡的存在呢? 关于超模博弈的另一个重要结论便与最小均衡和最大均衡的比较静态特征有关。为介绍这一结论, 我们看定义在 $\mathbb{R}$ 上的函数 $f(x)$ 。

图 5.8 中给出的递增函数有四个不动点。虚线 $f'(x)$ 是 $f(x)$ 的向上平移。在函数 $f(x)$ 向上移动时, 最大不动点和最小不动点向右移动(变大)。还有一些不动

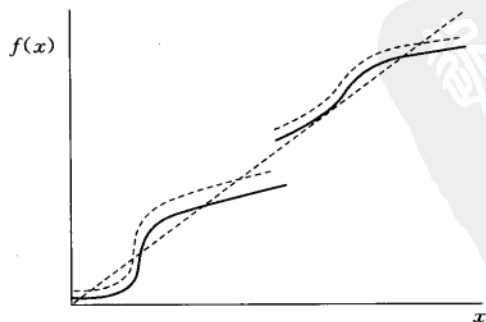


图 5.8 极值不动点的比较静态性

点则向左移动。最小不动点和最大不动点一般是函数从上方穿越 45° 线产生的结果,它们的性质类似。这一特点构成了比较静态分析的基础。

定理 5.16 在一般超模博弈中,最小均衡和最大均衡在 p 上递增。

图 5.8 指出,其他的均衡可能并不具有这一特征。Echenique(2002)指出,在具有互补性的博弈中,不具有最大均衡和最小均衡所具有的比较静态特征的一切均衡点都是不稳定的。换句话说,我们有明确的理由探讨具有单调比较静态特征的均衡点。

回到 Kahn 和 Kenney 的竞争性假说上,候选人 1 最大化 $p(s_1, s_2, \gamma) - c_1 s_1$, 候选人 2 最大化 $1 - p(s_1, s_2, \gamma) - c_2 s_2$, 其中 $p(s_1, s_2, \gamma)$ 决定于外生的竞争程度 $\gamma \in \mathbb{R}$ 。如果竞选支出影响竞争程度,并且在竞争更激烈的竞选中,选民们的关注程度更高,则当 s_2 或 γ 更高时, s_1 对 $p(s_1, s_2, \gamma)$ 的增量影响就更高。在 s_2 的增量效应上,也有同样的结论。就是说,收益函数是超模的假设与 Kahn 和 Kenney 的论点相一致。如果选择集为紧集,且 $p(s_1, s_2, \gamma)$ 在 s_1 上具有上半连续性,且 $1 - p(s_1, s_2, \gamma)$ 在 s_2 上具有上半连续性,我们就能够(利用定理 5.16 和定理 5.17)断定,均衡存在。而且,在最大均衡和最小均衡中, s_i 在 γ 上递增。这是个很强的结论;在 $p(\cdot, \cdot, \cdot)$ 的各种具体形式下,这一结论始终成立。单调比较静态的威力在于,它能使我们找到产生某一结果的最小结构。^①

5.12 精炼 Nash 均衡

在第 5.2.2 节的多数通过规则下的投票博弈中,任何策略组合,只要其中没有人是关键投票人,就是 Nash 均衡。我们不需要考虑这些均衡,因为它们包含弱劣策略。但是,在这一博弈中,剔除弱劣策略并不是我们精炼 Nash 均衡集的唯一方法。我们并不假定每个参与者都以概率 1 使用自己的最优反应,而假设每个参与者以某个很小的概率颤抖,就是说,以某个很小的概率使用别的策略。为使分析简单,我们看一个三人博弈。在这个博弈中,两个投票人偏好 D , 一个投票人偏好 R 。

假设每个参与者都以 ε 的概率使用每个纯策略,这里的 ε 是个小于 $\frac{1}{2}$ 的很小的正数。^②这一假设的含义是,投票人有可能犯错误,从而使得所有策略都被以某个最小概率在使用。

在这种情况下,最大化所偏好的候选人胜出的概率,显然是最优反应。这一目标和最小化最不喜欢的候选人胜出的概率是同一回事。所以,我们只需计算在各种策略组合下,每个候选人胜出的概率。候选人 R 胜出的概率为:

① 对超模博弈和单调比较静态等特征的深入讨论,可参见 Topkis(1998)。Ashworth 和 Bueno de Mesquita(2006)集概括了有关结论和它们在政治科学中的应用。

② 选民们可能不投票给他们本应支持的候选人的观点,在 2000 年美国总统选举后,再次具有了实质性意义。

$$\Pr(R \mid \text{三票支持 } R) = (1 - \varepsilon)^3 + 3(1 - \varepsilon)^2 \varepsilon$$

$$\Pr(R \mid \text{二票支持 } R) = (1 - \varepsilon)^3 + (1 - \varepsilon)^2 \varepsilon + (1 - \varepsilon) \varepsilon^2$$

$$\Pr(R \mid \text{一票支持 } R) = \varepsilon^2(1 - \varepsilon) + 2(1 - \varepsilon)^2 \varepsilon + \varepsilon^3$$

$$\Pr(R \mid \text{零票支持 } R) = \varepsilon^3 + 3\varepsilon^2(1 - \varepsilon)$$

不难证明, $\varepsilon < \frac{1}{2}$ 意味着共和党人(R)胜出的概率在支持他的选票数量上严格递增。

首先看在原始博弈中“差的”均衡,即R全票胜出的均衡。在存在颤抖的情况下,这一结果能生存下来吗?(就是说,每个选民都以 $1 - \varepsilon$ 的最大概率投票给R吗?)共和党选民通过以概率 $1 - \varepsilon$ 投票给R来最大化R胜出的概率。再看民主党选民的选择。他可以选择与这一均衡保持一致,即投票给R。此时,R胜出的概率是 $(1 - \varepsilon)^3 + 3(1 - \varepsilon)^2 \varepsilon$ 。他也可以偏离这一均衡而投票给D。此时,R胜出的概率是 $(1 - \varepsilon)^3 + (1 - \varepsilon)^2 \varepsilon + (1 - \varepsilon) \varepsilon^2$ 。他的偏离行为使得R胜出的概率下降了 $2(1 - \varepsilon)^2 \varepsilon - 2(1 - \varepsilon) \varepsilon^2 > 0$ 。显然,民主党选民会偏离这一均衡。同样地,所有选民都投票给D的均衡,也不会在这种颤抖中存活下来。

精炼均衡集的这种思想,关注的是在参与者微小失误时能生存下来的均衡。这一思想源于 Reinhardt Selten(1965),他称这种均衡为精炼均衡。精炼均衡的标准定义如下:

定义 5.20 给定 $\varepsilon > 0$, “ ε 约束下的”均衡是个完全混合的策略组合 σ^* , 使得对每个参与者 i , σ_i^* 是下面的最大化问题的解:

$$\max_{\sigma_i} (\sigma_i, \sigma_{-i}^*), \text{ s. t. 对每个 } s_i, \sigma_i(s_i) \geq \varepsilon.$$

精炼均衡是在 ε 趋向零时,由 ε 约束下的均衡构成的某个序列的极限。

不难看出,全票胜出的投票均衡并不满足这一定义。在 ε 约束下的均衡中,所有参与者以概率 ε 投票给他们最不喜欢的候选人。结果,在这些均衡所趋向的极限均衡上,投票给最不偏好的候选人的概率为0。实际上,唯一的精炼均衡是所有人投票给自己所支持的候选人。

精炼均衡有许多理想特征。首先,所有的精炼均衡都是原始博弈的不受 ε 约束的 Nash 均衡。因此,精炼均衡集是 Nash 均衡集的真子集。其次,利用前面证明 Nash 均衡存在性时使用的方法能够证明,有限标准式博弈至少有一个精炼均衡。^①

5.13 应用:私人提供公共品

自 Mancur Olson 的著作《集体行动的逻辑》(The Logic of Collective Action)

① ε 约束下的混合策略空间是紧的、凸的和非空的,这是整个证明的起点。当然了,在证明了 ε 约束下的均衡的存在性后,还必须证明它收敛于 Nash 均衡。

于1965年出版以来,在什么条件下,个人理性的决策主体愿意发生私人成本而为生产公共品提供资金的问题,成为政治科学中的一个核心问题。本节在 Palfrey 和 Rosenthal(1984)的基础上,介绍这样一种模型。

由 n 个人构成的团体,要决定是否生产某种公共品而由个人提供资金。生产这种公共品需要至少一个人提供资金。如果公共品被生产出来,每个人从中获得 1 个单位的效用。而任何提供资金的人,发生成本 $c < 1$ 。每个人的策略集为{提供资金,不提供资金}。我们把“提供资金”定义为 $s_i = 1$,把“不提供资金”定义为 $s_i = 0$ 。由于还要分析混合策略均衡,所以定义 σ_i 为第 i 个人提供资金的概率。每个人,如果提供了资金,得到的收益为 $1 - c$; 如果没有提供资金而其他人提供了资金,得到的收益为 1; 如果没有人提供资金,得到的收益为 0。这个模型包含有搭便车动机。每个人都希望其他人提供资金,这样他就可以得到好处而不必承担成本。

首先探讨纯策略均衡集。容易看出,这一博弈有纯策略 Nash 均衡,在这一均衡中,对每个 i , $s_i = 1$, $s_{-i} = 0$ 。在每个这样的均衡中,第 i 个人得到的收益为 $1 - c$, 其他人各得到 1 个单位的收益。第 i 个人不会偏离自己的均衡策略,因为如果偏离,他的收益就会下降到 0。其他人也不会偏离自己的均衡策略而提供资金,因为这会使自己的收益从 1 下降到 $1 - c$ 。其他任何策略组合都不是纯策略 Nash 均衡。首先,对所有的 i ,如果 $s_i = 0$,则任何人通过提供资金,都能提高自己的收益。所以,这种策略组合不可能是均衡。我们在某个策略组合中,不止一个人提供了资金,任何提供资金的人,单方面停止出资能增加自己的效用——公共品的生产可以由其他人提供的资金支持。

上述均衡与 Olson 的决策理论分析模型得到的结论完全不同。Olson 认为,搭便车行为导致公共品生产处于无效率状态。而在 Palfrey-Rosenthal 博弈中,所有的纯策略均衡都具有帕累托效率——公共品得以提供,而投入的资金保持在所需的最低水平上。但是,有理由怀疑这种纯策略均衡不是对这种博弈实际运行方式的真实描述。首先,有如此多的 Nash 均衡,参与者们如何协调各自的选择从而达到某个均衡? 其次,每个纯策略 Nash 均衡都要求事前完全相同的参与者在博弈中使用不同策略。相同的参与者使用相同的策略的均衡似乎更合理。所以,我们来探讨对称的混合策略均衡,在这种均衡中,对所有的 i , $\sigma_i = \sigma$ 。这一限制条件与我们对内生于纯策略 Nash 均衡中的不对称性的批判相一致。

前面说过,参与者 i ,只有在他对某个混合策略的支持集中的各个纯策略无差异时,才会使用这一混合策略。就是说,当且仅当 $u_i(\text{提供资金}, \sigma) = u_i(\text{不提供资金}, \sigma)$ 时,参与者 i 才会其他 $n - 1$ 个参与者都以概率 σ 提供资金时,使用策略 $0 < \sigma_i < 1$ 。至少有另外一人提供资金的概率是 $1 - (1 - \sigma)^{n-1}$, 所以无差异条件为:

$$1 - c = 1 - (1 - \sigma)^{n-1}$$

利用这一条件,可以求出提供资金的概率的均衡值是 $\sigma = 1 - c^{1/(n-1)}$ 。这一均衡与 Olson 的结论更一致。随着 n 变大, σ 趋于 0。但是,如果 c 趋于 0, σ 趋于 1。

多人出资

现在拓展这一模型,规定在 n 个人中至少 k 个人提供资金,这件公共品才能够生产出来。由于公共品的生产需要多个人提供资金,所以,所有人都不提供资金的策略组合是 Nash 均衡。有人如果偏离这一均衡而独自提供资金,就会一无所得且发生损失 c 。这一模型中存在着许多纯策略均衡,其中,正好有 k 个人提供资金。在这种均衡中,没有提供资金的人没有动机改变自己的选择(就是说,他没有动机提供资金);改变自己的选择,要发生成本 c ,却不改变获得公共品的概率。而提供资金的人,如果改变自己的选择,就会阻碍公共品的提供;他从节省下这笔资金中获得的价值,小于由于失去公共品而丧失的收益,所以,对提供资金的人来说,改变选择不是合理行为。所以,在这一模型中存在着一种均衡。在均衡中,总有由 k 个人组成的团体提供资金。根据组合数学,存在 $\binom{n}{k}$ 个这种均衡在其中, k 个人提供资金。^①

这些纯策略均衡,甚至比前述单人出资博弈中的纯策略均衡,更不合理。我们来计算这一博弈的对称混合策略均衡,用 σ 表示其中任何一个人提供资金的概率。由于其他 $N \setminus \{i\}$ 个参与者都在随机化自己的选择,我们用随机变量 x_{-i} 表示除 i 以外的所有人中出资人的数量。参与者 i 从提供资金中获得的收益是:

$$\Pr(x_{-i} < k-1) \cdot 0 + \Pr(x_{-i} \geq k-1) \cdot 1 - c$$

他从不提供资金的策略中获得的收益是:

$$\Pr(x_{-i} < k) \cdot 0 + \Pr(x_{-i} \geq k) \cdot 1$$

和前面一样,在混合策略均衡中,参与者在提供资金和不提供资金两个策略之间无差异。让上述两个收益相等,进行一番代数运算后,就可得到使用混合策略的必要条件:

$$\Pr(x_{-i} = k-1) = c$$

这一条件有相当好的直观解释。只有在正好有另外 $k-1$ 个人提供资金时,提供资金才给 i 带来正收益。因此,从为公共品提供资金中获得的收益,必须得乘以有一位出资人是关键出资人的概率。由于所有人都使用混合策略,所以这一预期收益必须等于出资成本 c 。

由于所有人独立地使用混合策略,我们可以算出 $\Pr(x_{-i} = k-1)$ 的值:在成功概率为 σ 时,它等于 $n-1$ 次试验中正好有 $k-1$ 次试验取得成功的概率。概率论提供了标准结论:

^① $\binom{n}{k} = \frac{n!}{k!(n-k)!}$ 被称为两项式系数,表示从 n 个物体中抽取 k 个所实现的组合的数量。

$$\Pr(x_{-i} = k-1) = \binom{n-1}{k-1} \sigma^{k-1} (1-\sigma)^{n-k}$$

求解这一对称混合策略均衡,就是要找到作为方程:

$$\binom{n-1}{k-1} \sigma^{k-1} (1-\sigma)^{n-k} = c$$

的解的 σ 值。

上式可以简化为:

$$\sigma^{k-1} (1-\sigma)^{n-k} = \frac{(k-1)!(n-k)!}{(n-1)!} c$$

在探讨解集的特征前,我们先看几个例子。首先,在 $n=5$ 和 $k=3$ 时,上述均衡条件为 $\sigma^2 (1-\sigma)^2 = c/6$ 。图 5.9 中的两条实线分别为这个式子等号两边的函数的图形。只要 c 足够低,这个方程就有两个解 σ_L^* 和 σ_H^* 且 $\sigma_L^* < \frac{1}{2} < \sigma_H^*$, 它们表示不同的混合策略均衡。均衡的比较静态性质决定于所探讨的均衡是 σ_L^* 还是 σ_H^* 。从图 5.9 上可以看出,作为 c 的函数,均衡混合策略是如何变化的。增加 c , 会提高 σ_L^* , 降低 σ_H^* 。

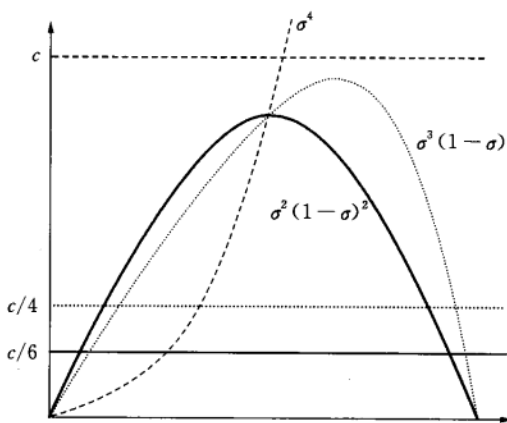


图 5.9 Palfrey-Rosenthal 出资博弈中的均衡

图 5.9 中还画出了 $k=4$ 和 $k=5$ 两种情况,它们所对应的均衡条件分别为 $\sigma^3(1-\sigma) = c/4$ 和 $\sigma^4 = c$ 。 $k=4$ 时的情况类似 $k=3$, 也有两个混合策略均衡。这里要注意的是,如果 $\sigma > \frac{1}{2}$, $\sigma^3(1-\sigma) > \sigma^2(1-\sigma)^2$ 。这一事实结合 $\frac{c}{4} > \frac{c}{6}$, 意味着 σ_H^* 在 k 上递增。如果 $\sigma < \frac{1}{2}$, 由于 $\sigma^3(1-\sigma) < \sigma^2(1-\sigma)^2$, σ_L^* 也在 k 上递增。在 $k=5$ 时, σ^4 在 $0 \leq \sigma \leq 1$ 上单调递增,这意味着只存在一个混合策略均

衡。和 $k = 4$ 时的 σ_L^* 相比,在这一均衡上,出资概率更高,并且它在 c 上递增。

上述例子中的许多结论具有普遍性。首先,无论 n 和 k 为多大,最多只能有两个混合策略均衡。设 $c(\sigma)$ 是使得 σ 为均衡混合策略的成本水平,换句话说,

$$c(\sigma) = \binom{n-1}{k-1} \sigma^{k-1} (1-\sigma)^{n-k}$$

对 σ 求导,得到:

$$c'(\sigma) = \binom{n-1}{k-1} [(k-1)(1-\sigma) - (n-k)\sigma] \sigma^{k-2} (1-\sigma)^{n-k-1}$$

如果 $k < n$, 函数 $c(\sigma)$ 是单峰的(在全局最大值的左边递增,在右边递减),因为:

$$\text{如果 } \sigma \begin{matrix} < \\ > \end{matrix} \frac{k-1}{n-1}, \text{ 则 } c' \begin{matrix} \geq \\ < \end{matrix} 0$$

如果 $1 < k < n$ 且 c 小于 $c \frac{k-1}{n-1} \equiv c_{\max}$, 则存在两个对称的混合策略均衡。如果 $c > c_{\max}$, 则没有均衡。^①如果 $k = n$ 或者 $k = 1$, 只要 $c < 1$, 就存在一个对称混合策略均衡。

在 $\sigma_H^* > \frac{k-1}{n-1}$ 的均衡中, 出资概率在 c 和 n 上递减, 在 k 上递增。在 $\sigma_L^* < \frac{k-1}{n-1}$ 的均衡中, 出资概率在 c 和 n 上递增。在 $\sigma_L^* < \frac{k-1}{n}$ 的均衡中, 出资概率在 k 上递减; 在 $\sigma_L^* \in \left(\frac{k-1}{n}, \frac{k-1}{n-1}\right)$ 的均衡中, 出资概率在 k 上递增。^②和 $k = 1$ 时的情况一样, 随着 n 上升, 出资概率趋于 0。^③

5.14 习题

习题 5.1 证明 Nash 均衡的两个定义等价。提示: 首先证明某个策略组合如果满足第一个定义, 则肯定满足第二个定义。然后证明它如果满足第二个定义, 则肯定满足第一个定义。

习题 5.2 证明在囚徒两难博弈, 没有非纯策略 Nash 均衡。

习题 5.3 和 5.4 在下面的标准式博弈中,

- ① 在 $c = c_{\max}$ 这一不大可能发生的事件中, 存在唯一一个对称混合策略均衡。
- ② Palfrey 和 Rosenthal(1988)的附录中, 给出了证明这些结论所需的一些推导。
- ③ 这最后一个结论的证明要用到大数法则。由于 $c(\sigma)$ 是两项分布的概率函数, 所以随着 n 变大, 除了在点 $\frac{k-1}{n-1}$ 以外, 它处处趋于 0。因此, σ_L^* 和 σ_H^* 肯定收敛于 $\frac{k-1}{n-1}$, 而 $\frac{k-1}{n-1}$ 本身收敛于 0。

1\2	<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
<i>w</i>	1, 2	1, 3	2, 1	5, 1
<i>x</i>	4, 2	2, 4	3, 4	4, 3
<i>y</i>	3, 2	2, 3	4, 5	2, 2
<i>z</i>	3, 2	1, 2	2, 2	1, 4

(1) 哪些策略在反复剔除严格劣策略中存活下来?

(2) 哪些策略组合是纯策略 Nash 均衡?

(3) 哪些策略组合是混合策略 Nash 均衡?

习题 5.5 找出下面的博弈中的混合策略 Nash 均衡。

1\2	<i>L</i>	<i>R</i>
<i>T</i>	2, 1	0, 2
<i>B</i>	1, 3	3, 0

习题 5.6 和 5.7 在下面的标准式博弈中,

1\2	<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
<i>u</i>	2, 0	3, 1	4, 0	6, 1
<i>m</i>	3, 4	4, 2	6, 3	4, 1
<i>d</i>	4, 4	3, 3	5, 6	6, 2

(1) 哪些策略是弱劣策略?

(2) 哪些策略在反复剔除严格劣策略中存活下来?

(3) 找出混合策略 Nash 均衡和/或纯策略 Nash 均衡。

习题 5.8 一个由三人组成的立法机构实行多数通过的投票法则。设 y 是当前实施的法案, x 是备选法案。游说者 L_1 希望通过法案 x , 游说者 L_2 希望依然实行现行法案 y 。每个游说者可以向立法人员“行贿”, 求他投票支持自己所偏好的法案。假设贿赂只能为 p ($p > 0$ 是个常数)。一位立法人员如果只得到 L_1 提供的贿赂, 就投票给 x ; 如果只得到 L_2 提供的贿赂, 就投票给 y ; 如果得到双方的贿赂或者没有得到任何贿赂, 就投票给 y 。下面的效用函数反映这两个游说者的偏好:

$$U_{L_1}(x) = u^* - pB_{L_1}$$

$$U_{L_1}(y) = -pB_{L_1}$$

$$U_{L_2}(x) = -pB_{L_2}$$

$$U_{L_2}(y) = u^* - pB_{L_2}$$

其中的 B_{L_1} 是 L_1 行贿的立法人员数量(其值为 0, 1, 2 或 3), B_{L_2} 是 L_2 所行贿的立

法人员数量(其值为 0, 1, 2 或 3)。设 $u^* \geq 3p$ 。每个游说者的一个纯策略, 规定了行贿对象(如果他行贿的话)。例如, $(b, 0, b)$ 表示向立法人员 1 和立法人员 3 行贿但不给立法人员 2 贿赂。一个混合策略则是八个可能的纯策略上的一个概率分布。

(1) 探讨这一博弈的(混合策略或纯策略)Nash 均衡的特征。提示: 首先剔除严格劣策略。

(2) 在 Nash 均衡中, 法案 x 得以通过的概率是多少? 这一概率如何对参数 p 和 u^* 的变化做出反应?

习题 5.9 说两个候选人 a 和 b 竞选公职。政策空间由三个点组成, 为 $\{-1, 0, 1\}$; 选民们有对称的伞状效用函数。候选人不知道中间选民的位置, 他们的推断是: 中间选民位于点 -1 的概率是 π 且 $\pi < \frac{1}{2}$, 位于点 1 的概率也是 π , 位于点 0 的概率是 $1 - 2\pi$ 。设候选人 a 比候选人 b 更优秀(她的履历更杰出、更有吸引力或者所在政党更受欢迎)。设两人之间的这种质量差异很小。如果 a 和 b 到中间选民的距离相等(或者 a 更接近中间选民), 则 a 胜出; 但是, 如果 b 更接近中间选民, 则 b 胜出。在这个标准式博弈中, 两个候选人同时在政策空间 $\{-1, 0, 1\}$ 中选择一个政策。

(1) 在这一博弈中, 存在纯策略 Nash 均衡吗? 如果存在, 请分析它们的特征。

(2) 在这一博弈中, 存在混合策略 Nash 均衡吗? 如果存在, 请分析它们的特征。

习题 5.10 在 Downsian/Hotelling 模型中, 候选人有各自的政策偏好, 中间选民的理想点在 $[0, 1]$ 上均匀分布。设一位候选人的理想点在 $[0, \frac{1}{3}]$ 上均匀分布, 另一位候选人的理想点在 $[\frac{1}{2}, 1]$ 上均匀分布。分析这一博弈中 Nash 均衡的特征。

习题 5.11 在各种 Hotelling 模型里, 证明如果有三个参与者, 不存在纯策略均衡。如果有四个参与者且每个参与者最大化自己的得票份额, Nash 均衡是什么?

习题 5.12 证明定理 5.5。

习题 5.13 在国际外部性博弈中, 分析 $k_1 = k_2$ 时的纯策略 Nash 均衡。

习题 5.14 (*) 证明定理 5.6。

习题 5.15 (*) 证明: 上半连续的对应, 如果是单值的, 则为连续函数。

习题 5.16 设 s_i^* 是一阶条件 $\beta - 2ks_i = 0$ 的解。用隐函数定理, 求 $\partial s_i^* / \partial \beta$ 和 $\partial s_i^* / \partial k$ 。

习题 5.17 利用方程组 5.1, 求出 $b_2(\cdot, \cdot; \cdot)$ 的导数。

习题 5.18 利用方程组 5.1, 求最优反应的二阶(和交叉)偏导数。

习题 5.19 证明: $f(\cdot, \cdot)$ 如果在 (x, p) 上有递增差, 则对所有的 $p \leq p'$ 和

$$x \leq x', f(x, p') - f(x, p) \geq f(x', p') - f(x', p).$$

习题 5.20(*) 证明定理 5.11 中的(1)~(3)部分。

习题 5.21(*) 设 X 是 $\mathbb{R}^n$ 的紧子集, R 是 X 上的下连续偏序。证明: 如果 $u(x)$ 是表示 X 上的偏序 R 的效用函数, 则 $u(\cdot)$ 在 X 上是上半连续的。如果 $u(x)$ 在 X 上是上半连续的, 则它所表示的偏好关系是下连续的。

习题 5.22(*) 证明: 每个有限博弈有至少一个精炼均衡。

习题 5.23 在 Palfrey-Rosenthal 的公共品出资博弈中, 构造不对称 Nash 均衡: l 个参与者出资 ($\sigma_i = 1$), m 个参与者不出资 ($\sigma_i = 0$), 其余的 $n - m - l$ 个参与者使用混合策略 $\sigma_i = q \in (0, 1)$ 。证明: 如果 $l > 0$ 或 $m > 0$, 则 q^* 是唯一的。

习题 5.24 在 Palfrey-Rosenthal 模型中, 如果公共品最终没有生产出来(因为出资不足), 则各个人所提供的资金被全额退回。探讨这一博弈中的纯策略均衡和混合策略均衡。

解
答
PDG

▶ 6

标准式贝叶斯博弈

前一章的标准式博弈假定参与者拥有完全信息；或者在存在不确定性时，拥有相同的概率推断。但是，这种假设有时并不合理。竞选公职的候选人比选民

更了解自己的政策偏好；利益集团可能比立法人员更了解政策和结果之间的关系；一个国家可能比对手更了解自己的军事力量。在许多情况中，如果忽视信息不对称性，就会漏掉许多重要的策略性动机。我

表 6.1 抓捕恐怖分子

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	2, 2	0, 1
抓捕行动人员	1, 0	1, 1

们回到上一章抓捕恐怖分子博弈，表 6.1 给出了这一博弈的结构。

CIA 知道，FBI 偏好抓获行动人员而不是一无所获。FBI 知道 CIA 知道这一点，如此等等。但是，如果 CIA 并不确定 FBI 是偏好抓捕行动人员还是偏好一无所获——FBI 可能认为，抓到行动人员所实现的国土安全收益不足以补偿抓捕成本——这一博弈就会发生重大变化。CIA 如果意识到这一点，就会知道实际进行的可能是表 6.2 中的博弈。

表 6.2 抓捕恐怖分子

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	2, 2	0, 1
抓捕行动人员	0, 0	0, 1

CIA 甚至会认为，FBI 还可能有另一种偏好：抓捕行动人员比抓捕核心分子更有价值。因为一般情况下，行动人员在压力下会坦白而核心分子会抗拒到底。此时，CIA 知道自己面临的可能是表 6.3 中的博弈。

如果无法确切地知道 FBI 如何评价各个结果，CIA 很难确定自己应该使用哪一策略。在表 6.1 的博弈中，纯策略 Nash 均衡是（抓捕核心分子，抓捕核心分子）

和（抓捕行动人员，抓捕行动人员）。但是，如果 CIA 认为 FBI 在进行表 6.2 或表 6.3 中的博弈，我们就不清楚 CIA 偏好哪一策略，我们也就无法求解 Nash 均衡或者剔除劣策略。在构造模型时，我们即使知道 FBI

表 6.3 抓捕恐怖分子

FBI\CIA	抓捕核心分子	抓捕行动人员
抓捕核心分子	1, 2	0, 1
抓捕行动人员	2, 0	2, 1

的收益,并且知道哪个策略对 FBI 是最优反应或者是个劣策略,依然远远不够,因为这些解概念要求 CIA 清楚地知道 FBI 的收益。

在本节所介绍的模型中,参与者不知道对手的收益,这一特点被称为**不完全信息**。我们介绍分析此类模型所需的工具。标准的处理方法(源于 Harsanyi, 1967~1968)是把这种博弈转换为另一种博弈:一个虚构的参与者(通常被称为“自然”)首先行动,从某个给定的概率分布中抽取各个参与者的效用函数。参与者都知道这一概率分布,但无法观察到自然抽取的结果。一般情况下,参与者能观察到抽取结果中的某些方面,如自己的收益。在自然行动之后,参与者同时选择自己的行动。这种博弈有**完全的但非完美的信息**。

在抓捕恐怖分子博弈中,自然从已知概率分布中随机选择 FBI 的偏好(或类型)。FBI 知道自己的类型,选择自己的行动。CIA 在选择自己的行动时,不知道 FBI 的类型,但知道每种类型的发生概率。CIA 的策略有些复杂。在选择自己的策略时,CIA 必须推断每一种 FBI 类型所使用的策略。根据所形成的推断,CIA 比较自己从“抓捕核心分子”策略中得到的期望效用和从“抓捕行动人员”策略中得到的期望效用。另一方面,每种 FBI 类型,根据自己对 CIA 策略的推断,选择自己的最优反应。本质上说,Harsanyi 的方法是把 CIA 不知道 FBI 的偏好的博弈,转换为一个新的博弈,在这个新博弈中,CIA 与从已知分布中抽取的三种 FBI 类型之一博弈。在新博弈中,每个参与者类型(即 CIA 和三种类型的 FBI)选择自己的策略。

这里需要注意公共知识(common knowledge)在不完全信息博弈中的重要性。本章开头说过,完全信息博弈假定博弈的所有构成因素——参与者、策略和收益——被所有参与者知道,并且所有参与者都知道这一点,如此等等。不完全信息博弈依然保留公共知识假设:所有参与者知道自然在选择参与者类型时所使用的概率分布。而且,所有参与者都知道,在自然所采取的行动中,有多少信息揭示给了对手。所有参与者都知道所有参与者知道这些细节,如此等等。

6.1 正式定义

我们现在调整基本的标准式结构 Γ , 引入关于参与者类型的非完美信息。^①在博弈 Γ 中,我们加入参与者类型、状态变量和这些随机变量上的概率分布。

(1) 类型:每个参与者 $i \in N$, 有一个由所有可能的类型构成的有限集 Θ_i 。 $\theta_i \in \Theta_i$ 为参与者 i 的一个类型; $\theta \in \Theta = \times_{i \in N} \Theta_i$ 是由 n 个参与者的类型构成的一个类型组合。 θ_{-i} 和 Θ_{-i} 分别表示除 i 以外的所有参与者的类型构成的一个组合和由所有这种类型组合 θ_{-i} 构成的集合。在前面的例子中,CIA 的类型集为一个单

^① 在学习本节前,你可能需要重温数学附录中的概率论部分。

元素集合, FBI 的类型集中有三个元素。

(2) 随机状态变量: 博弈 Γ 中可能包含一个随机状态变量 $\omega \in \Omega$, 其中 Ω 是有限集。^①

(3) 自然随机地做出选择: 在博弈开始时, 自然从一个联合概率分布中, 选择参与者类型向量 $\theta = (\theta_1, \dots, \theta_n) \in \Theta = \prod_{i \in N} \Theta_i$ 和 $\omega \in \Omega$ 。在这个联合分布中, (θ, ω) 的发生概率是 $p(\theta, \omega)$ 。函数 $p(\theta_{-i}, \omega | \theta_i)$ 是在给定 θ_i 时, (θ_{-i}, ω) 的条件概率。

(4) 策略: 每个参与者从策略集 S_i 中选择行动 s_i 。

(5) 期望效用: 对每个可能的策略组合 s 、类型组合 θ 和状态 ω , 参与者 i 的效用函数是 $u_i(s_i, s_{-i}, \theta, \omega)$ 。给定类型 θ_i , 参与者 i 从策略组合 s 中得到的条件期望效用是:

$$EU_i(s; \theta_i) = \sum_{\theta_{-i} \in \Theta_{-i}} \sum_{\omega \in \Omega} p(\theta_{-i}, \omega | \theta_i) u_i(s, \theta_i, \theta_{-i}, \omega)$$

于是, 标准式贝叶斯博弈为 $\langle N, \Omega, \{S_i, \Theta_i, u_i(\cdot, \dots, \cdot)\}_{i \in N}, p(\cdot, \cdot) \rangle$ 或者缩写为 $\langle N, \Omega, S, \Theta, u, p \rangle$ 。正如标准式博弈可以有无限策略空间, 贝叶斯博弈也可以有无限类型空间和无限行动空间。^②后文将给出一些这方面的例子。

在贝叶斯博弈的策略组合中, 必须包含每个参与者类型所使用的策略。因此, 参与者 i 的一个策略便是一个函数 $\phi_i: \Theta_i \rightarrow S_i$, 它为每个可能的类型 $\theta_i \in \Theta_i$, 规定了一个策略 $s_i \in S_i$ 。在引入了不确定性之后的抓捕恐怖分子博弈中, FBI 的类型集为 $\Theta_{FBI} = \{\text{标准型}, \text{抓捕核心分子型}, \text{抓捕行动人员型}\}$ 。 $\phi_{FBI}(\text{标准型}) = \text{抓捕核心分子}$ 、 $\phi_{FBI}(\text{抓捕核心分子型}) = \text{抓捕核心分子}$ 和 $\phi_{FBI}(\text{抓捕行动人员型}) = \text{抓捕行动人员}$, 是 FBI 的一个策略。

在贝叶斯标准式博弈中, 与 Nash 均衡相对应的概念是贝叶斯 Nash 均衡。由于策略是类型的函数, 对最优反应的推导有些复杂。所以, 我们利用 Nash 均衡的第二个定义来定义贝叶斯 Nash 均衡。

定义 6.1 在标准式贝叶斯博弈 $\langle N, \Omega, S, \Theta, u, p \rangle$ 中, 一个贝叶斯 Nash 均衡是一个策略组合 $(\phi_1^*(\cdot), \dots, \phi_n^*(\cdot))$, 使得对每个 $i \in N$, 每个 $s'_i \in S_i$, 和每个 $\theta_i \in \Theta_i$,

$$EU_i(\phi_i^*(\theta_i), \phi_{-i}^*(\cdot); \theta_i) \geq EU_i(s'_i, \phi_{-i}^*(\cdot); \theta_i) \quad (6.1)$$

就是说, 在贝叶斯 Nash 均衡中, 给定所有其他参与者类型所使用的策略和类型上的概率分布, 每个参与者类型选择最大化自己的期望效用的策略。

① 严格地说, 随机变量 ω 没有必要出现在贝叶斯博弈的定义中。我们可以定义仅决定于 θ 的收益上的期望效用, 并假定参与者最大化期望效用。在这里, 通过引入状态向量, 我们实际上定义了 θ 和 ω 上的效用。前一章指出, 在标准式博弈中, 收益可以是随机的。因此, 这两种处理方法等价。我们明确引入随机状态, 是为了说明不确定性进入博弈的两种方式(关于其他人的知识上的不确定性, 即 θ_i ; 和关于行动如何影响收益上的不确定性, 即 ω)。

② 如果类型空间和行动空间不是有限空间, 在分析均衡时所用到的数学工具更加复杂。

现在对多类型抓捕恐怖分子博弈求解贝叶斯 Nash 均衡。为简化分析,假定三种 FBI 类型有相同的发生概率。设 FBI 的策略是: ϕ_{FBI} (标准型) = 抓捕核心分子, ϕ_{FBI} (抓捕核心分子型) = 抓捕核心分子, ϕ_{FBI} (抓捕行动人员型) = 抓捕行动人员。在这一博弈中, CIA 只有一种类型, 并且除 FBI 的类型外, 不存在其他不确定性。所以, 我们不必考虑 ω 。给定 FBI 的这个策略, CIA 的期望效用为

$$EU_{CIA}(\text{抓捕核心分子}, \phi_{FBI}(\cdot)) = \frac{2}{3}2 + \frac{1}{3}0 = \frac{4}{3}$$

$$EU_{CIA}(\text{抓捕行动人员}, \phi_{FBI}(\cdot)) = \frac{2}{3}0 + \frac{1}{3}1 = \frac{1}{3}$$

因此, CIA 对 $\phi_{FBI}(\cdot)$ 的最优反应是抓捕核心分子。现在检验是否有某一类型的 FBI 会偏离上述策略。对标准型 FBI 来说, 选择和 CIA 相同的策略是它的最优反应, 即“ ϕ_{FBI} (标准型) = 抓捕核心分子”是最优反应; 在 CIA 选择抓捕核心分子时, 对抓捕核心分子型 FBI 来说, 选择“抓捕核心分子”能带来最高收益, 所以, 它没有动机偏离这一选择; 在 CIA 抓捕核心分子时, 对抓捕行动人员型 FBI 来说, 选择“抓捕行动人员”能产生 2 个单位的效用, 选择“抓捕核心人员”能产生 1 个单位的效用。因此, 对 FBI 来说, 任何偏离行为都无利可图。所以, 上述策略组合构成贝叶斯 Nash 均衡。

6.2 应用: 贸易保护

两个国家考虑实施限制性贸易政策。设 $N = \{1, 2\}$; 每个国家有两种可能类型, 类型集为 $\Theta_i = \{u, b\}$ 。类型为 u 的国家想单方面限制从另一国家进口; 如果另一国家限制进口, 则类型为 b 的国家推行限制贸易的双边政策。两个国家的类型是独立抽取的, 类型为 u 的概率都是 $p \in (0, 1)$ 。每个国家的策略空间是 $S = \{l, f\}$, 其中 l 表示进口限制, f 表示自由贸易。国家 i 的收益是:

$$u_i(s_i, s_{-i}; \theta_i) = \begin{cases} 3, & \text{如果 } s_i = l, s_{-i} = f \text{ 且 } \theta_i = u \\ 2, & \text{如果 } s_i = f, s_{-i} = f \text{ 且 } \theta_i = u \\ 1, & \text{如果 } s_i = l, s_{-i} = l \text{ 且 } \theta_i = u \\ 0, & \text{如果 } s_i = f, s_{-i} = l \text{ 且 } \theta_i = u \\ 3, & \text{如果 } s_i = f, s_{-i} = f \text{ 且 } \theta_i = b \\ 2, & \text{如果 } s_i = l, s_{-i} = f \text{ 且 } \theta_i = b \\ 1, & \text{如果 } s_i = l, s_{-i} = l \text{ 且 } \theta_i = b \\ 0, & \text{如果 } s_i = f, s_{-i} = l \text{ 且 } \theta_i = b \end{cases}$$

在这一博弈中, 一个策略是一个映射 $s_i(\theta_i): \{u, b\} \rightarrow \{l, f\}$ 。这一博弈的一个特征是, 无论另一国家选择什么行动, 类型为 u 的国家总能从 l 中得到更高收益。如果

两个国家都为类型 u , 且这一事实是公共知识, 这一博弈便是囚徒两难: 每个国家都有占优策略, 那就是选择 l 。如果两个国家都为类型 b , 且这一事实是公共知识, 这一博弈就有两个纯策略 Nash 均衡 (f, f) 和 (l, l) 。

为求解贝叶斯 Nash 均衡, 我们首先猜测均衡策略, 然后验证它们是否满足均衡条件。类型为 u 的国家有占优策略 l , 所以, 在每个均衡中, 必然有 $s_i(u) = l$ 。因此, 可能的对称纯策略均衡为 $(s_i(u), s_i(b)) = (l, l)$ 或 $(s_i(u), s_i(b)) = (l, f)$ 。我们首先分析 $s_i(u) = l$ 和 $s_i(b) = f$ 构成均衡的可能性, 其中 $i = 1, 2$ 。如果国家 2 使用这一策略, 则 $s_2 = l$ 概率为 p , $s_2 = f$ 的概率为 $1 - p$ 。于是, 国家 1 在类型为 u 和 b 时的期望效用分别为:

$$EU_1(s_1, \theta_1 = u) = \begin{cases} p + (1-p)3, & \text{如果 } s_1 = l \\ 2(1-p), & \text{如果 } s_1 = f \end{cases}$$

$$EU_1(s_1, \theta_1 = b) = \begin{cases} 2(1-p) + p, & \text{如果 } s_1 = l \\ 3(1-p), & \text{如果 } s_1 = f \end{cases}$$

那么, 我们所猜测的策略是最优反应吗? 策略 $s_i(u) = l$ 是最优反应, 因为实施贸易壁垒是类型为 u 的国家的占优策略。如果 $3(1-p) \geq 2(1-p) + p$, 则 $s_i(b) = f$ 也是最优反应。而只要 $p \leq \frac{1}{2}$, 这一条件就成立。同样的结论也适用于国家 2。所以在 $p \leq \frac{1}{2}$ 时, 策略组合 $s_i(u) = l$ 和 $s_i(b) = f$ 构成贝叶斯 Nash 均衡。

现在来验证对这两个国家, $s_i(b) = l$ 是否是最优反应, 和在什么条件下是最优反应。如果在 θ_2 的任何值上, 国家 2 始终使用策略 $s_2(\theta_2) = l$, 则类型为 b 的国家 1 的期望效用函数为

$$EU_1(s_1, \theta_1 = b) = \begin{cases} 1, & \text{如果 } s_1 = l \\ 0, & \text{如果 } s_1 = f \end{cases}$$

不管 p 的值为多大, 策略组合 $(s_i(b) = l, s_i(u) = l)$ 构成贝叶斯 Nash 均衡。这就是说, 在这一博弈中, 始终存在一个贝叶斯 Nash 均衡, 两个国家都实行进口限制政策 (l, l) 。而且, 如果 $p \leq \frac{1}{2}$, 还存在另一个均衡: $(s_i(u) = l, s_i(b) = f)$ 。在这个均衡中, 如果两个国家都为类型 b , 就有自由贸易出现。这一结果的先验概率为 $(1-p)^2$ 。所以, 当每个国家都是双边类型且都认为其对手可能为双边类型时, 双边自由贸易政策就会出现。

6.3 应用: 陪审团投票

设有三个陪审员, 即 $N = \{1, 2, 3\}$, 负责宣判被告有罪或无罪。^①他们得集体

① 这是 Austen-Smith 和 Banks(1996)的陪审团模型的简化。

选出一个结果 $x \in \{c, a\}$, c 表示“宣判被告有罪”, a 表示“宣判被告无罪”。三个陪审员同时投票 $v_i \in S_i = \{c, a\}$, 在多数通过规则下生成最终判决结果。每个陪审员都不确定被告是有罪还是无罪。设 G 表示有罪, I 表示无罪, 状态变量集于是为 $\Omega = \{G, I\}$ 。每个陪审员对状态 G 有先验概率 $\pi > \frac{1}{2}$ 。如果被告有罪, 每个陪审员从有罪宣判中得到 1 个单位的效用, 从无罪宣判中得到 0 个单位的效用。如果被告无罪, 每个陪审员从无罪宣判中得到 1 个单位的效用, 从有罪宣判中得到 0 个单位的效用。

如果没有更多信息, 每个陪审员从有罪判决中得到 π 个单位的期望效用, 从无罪判决中得到 $1 - \pi$ 个单位的期望效用。由于 $\pi > \frac{1}{2}$, 所以在剔除弱劣策略中存活下来的 Nash 均衡中, 每个陪审员都投票认定被告有罪。

现在假设在投票之前, 每个陪审员都得到关于被告是否有罪的私人信号 $\theta_i \in \{0, 1\}$ 。这一信号包含有用信息: 和被告无罪时相比, 在被告有罪时, 每个陪审员更可能收到信号 $\theta_i = 1$ 。为简化分析, 假设在被告有罪时陪审员收到有罪信号 $\theta_i = 1$ 的概率, 和在被告无罪时收到无罪信号 $\theta_i = 0$ 的概率相同, 就是说, $\Pr(\theta_i = 1 | \omega = G) = \Pr(\theta_i = 0 | \omega = I) = p > \frac{1}{2}$, $\Pr(\theta_i = 0 | \omega = G) = \Pr(\theta_i = 1 | \omega = I) = 1 - p$ 。

在收到信号后, 陪审员 i 通过选择自己的投票策略 $v(\theta_i)$, 最大化做出正确决策的概率, 即不放过一个罪犯且不冤枉一个无辜者的概率。假设每个陪审员都使用诚实策略(sincere strategy): $v_i(1) = c$, $v_i(0) = a$ 。诚实策略要求在收到被告有罪信号后投票认定被告有罪, 在收到被告无罪信号后投票认定被告无罪。诚实策略构成贝叶斯 Nash 均衡——陪审员 1 愿意使用这一策略, 当他认为陪审员 2 和 3 使用诚实策略时。根据这些推断, 陪审员 1 投票认定被告有罪而得到的期望效用为:

$$\begin{aligned} & \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = G | \theta_1) \\ & + \Pr(\theta_3 = 1 \text{ 且 } \theta_2 = 0 \text{ 且 } \omega = G | \theta_1) \\ & + \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 1 \text{ 且 } \omega = G | \theta_1) \\ & + \Pr(\theta_2 = 0 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = I | \theta_1) \end{aligned}$$

陪审员 1 投票认定被告无罪而得到的期望效用为:

$$\begin{aligned} & \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = I | \theta_1) \\ & + \Pr(\theta_3 = 1 \text{ 且 } \theta_2 = 0 \text{ 且 } \omega = I | \theta_1) \\ & + \Pr(\theta_2 = 0 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = I | \theta_1) \\ & + \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 1 \text{ 且 } \omega = G | \theta_1) \end{aligned}$$

在上述两个期望效用表达式中, 最后两项相同, 所以, 在比较效用时, 这两项抵消掉。于是, 当且仅达

$$\begin{aligned} & \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = G | \theta_1) + \Pr(\theta_3 = 1 \text{ 且 } \theta_2 = 0 \text{ 且 } \omega = G | \theta_1) \\ & \geq \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = I | \theta_1) + \Pr(\theta_3 = 1 \text{ 且 } \theta_2 = 0 \text{ 且 } \omega = I | \theta_1) \end{aligned}$$

投票认定被告有罪为最优反应。由于状态变量和其他陪审员收到的信号的条件概率决定这些表达式的值,所以,陪审员 1 使用贝叶斯法则计算其中各项的值。假设陪审员 1 得到的信号为 $\theta_1 = 1$ 。根据贝叶斯法则:

$$\begin{aligned} & \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = G | \theta_1 = 1) \\ & = \Pr(\theta_3 = 1 \text{ 且 } \theta_2 = 0 \text{ 且 } \omega = G | \theta_1 = 1) \\ & = \frac{\pi p^2(1-p)}{\pi p + (1-\pi)(1-p)} \end{aligned}$$

$$\begin{aligned} & \Pr(\theta_2 = 1 \text{ 且 } \theta_3 = 0 \text{ 且 } \omega = I | \theta_1 = 1) \\ & = \Pr(\theta_3 = 1 \text{ 且 } \theta_2 = 0 \text{ 且 } \omega = I | \theta_1 = 1) \\ & = \frac{(1-\pi)p(1-p)^2}{\pi p + (1-\pi)(1-p)} \end{aligned}$$

于是,对陪审员 1 来说,如果:

$$2 \frac{\pi p^2(1-p)}{\pi p + (1-\pi)(1-p)} \geq 2 \frac{(1-\pi)p(1-p)^2}{\pi p + (1-\pi)(1-p)}$$

$v_1(1) = c$ 就是最优的。简化和调整这一不等式,得到:

$$\frac{\pi p^2(1-p)}{\pi p^2(1-p) + (1-\pi)p(1-p)^2} \geq \frac{1}{2}$$

在不等号左边,是在有两个信号为 $\theta = 1$ 和一个信号为 $\theta = 0$ 时,被告有罪的条件概率。换句话说,给定自己的信号和自己是关键投票人的推断,陪审员 1 如果认为被告更可能有罪而不是无罪,就会投票认定被告有罪。

在得到的信号为 0 时,陪审员 1 投票认定被告无罪的条件是:

$$\frac{\pi p(1-p)^2}{\pi p(1-p)^2 + (1-\pi)p^2(1-p)} \leq \frac{1}{2}$$

总之,在投票选择与所得到的私人信号相一致的任何贝叶斯均衡中,以下结论肯定成立:(1)如果 i 是关键投票人且 $\theta_i = 1$,则被告有罪的后验概率大于 $\frac{1}{2}$;

(2)如果 i 是关键投票人且 $\theta_i = 0$,则被告有罪的后验概率小于 $\frac{1}{2}$ 。

但是,Austen-Smith 和 Banks(1996)指出,在许多情况中,诚实策略与均衡行为并不一致。在参数 π 和 p 的某些值上,有的必要条件并不成立。陪审员可能选择其他策略。他们在收到某些信号时,可能会随机化自己的选择,或者不管得到什么信号都以相同方式投票,或者采用不同于其他陪审员的策略。Feddersen 和 Pessendorfer(1998)探讨了在投票规则和陪审员数量改变时,这一博弈的均衡

所具有的特征。^①

6.4 应用：在信号集为连续统时陪审团的投票行为

假设各个陪审员收到的不是两元信号，而是信号 $\theta_i \in [0, 1]$ ，其中 θ_i 产生自条件分布 $F(\theta_i|\omega)$ 。设这一分布函数有可微的密度函数 $f(\theta_i|\omega)$ ，并且 $f(\theta_i|\omega)$ 满足单调似然率条件。^①

定义 6.2 条件密度函数满足严格单调似然率条件 (SMLR)，如果在 $[0, 1]$ 上， $f(\theta_i|G)/f(\theta_i|I)$ 是 θ_i 的严格单调函数。

这一条件很重要。根据贝叶斯法则，

$$\Pr(G|\theta_i) = \frac{f(\theta_i|G)\pi}{f(\theta_i|G)\pi + f(\theta_i|I)(1-\pi)} = \frac{\frac{f(\theta_i|G)}{f(\theta_i|I)}\pi}{\frac{f(\theta_i|G)}{f(\theta_i|I)}\pi + (1-\pi)}$$

因此， $\Pr(G|\theta_i)$ 在 θ_i 上递增当且仅当 $f(\theta_i|G)/f(\theta_i|I)$ 在 θ_i 上递增。这就是说，SMLR 条件意味着在信号值更大时， $\omega = G$ 的后验概率更高。

为使分析简单，我们只探讨对称策略，即收到相同信号的陪审员选择相同策略。对称策略组合因此为映射 $v_i(\theta_i): [0, 1] \rightarrow \{c, a\}$ 。和两元信号模型一样，贝叶斯 Nash 均衡策略是在每个人都根据自己的私人信息和自己是关键投票人的推断而选择行动时的最优策略。一位陪审员，如果认为被告有罪的概率不小于 $\frac{1}{2}$ ，就投票认定被告有罪；如果认为被告有罪的概率不超过 $\frac{1}{2}$ ，就投票认定被告无罪。由于越大的

信号值越好地反映被告有罪，所以，陪审员的这种投票策略肯定是弱递增的。陪审员在 θ_i 的值小时，投票认定被告无罪；在 θ_i 的值高时，投票认定被告有罪。这种单调策略可以用折点 $\hat{\theta} \in [0, 1]$ 来定义。假设陪审员 $j \in N \setminus \{i\}$ 使用如下单调策略：

$$v_j(\theta_j) = \begin{cases} c, & \text{如果 } \theta_j \geq \hat{\theta} \\ a, & \text{如果 } \theta_j < \hat{\theta} \end{cases}$$

如果除 i 以外的所有陪审员都使用这种折点策略，则给定信号 θ_i 和 i 是关键投票人的推断， $\omega = G$ 的后验概率为：

$$\Pr(G | p_{iv}, \theta_i; \hat{\theta}) = \frac{\pi f_G \hat{F}_G^{\pi-1} [1 - \hat{F}_G]^{1-\pi}}{\pi f_G \hat{F}_G^{\pi-1} [1 - \hat{F}_G]^{1-\pi} + (1-\pi) f_I \hat{F}_I^{\pi-1} [1 - \hat{F}_I]^{1-\pi}}, \quad (6.2)$$

^① Duggan 和 Martinelli(2001)和 Meirowitz(2002)探讨了这种情况。

其中, $f_{\omega} = f(\theta_i | \omega)$, $\hat{F}_{\omega} = F(\hat{\theta} | \omega)$ 。我们把这一表达式的推导留做习题。这一概率是参数 $\hat{\theta}$ 的函数。在这一模型中, 陪审员使用折点策略的对称均衡的存在, 取决于满足下述关系的 $\hat{\theta}$ 的值:

$$\Pr(G | piv, \hat{\theta}; \hat{\theta}) = \frac{1}{2}$$

且在 $\theta_i < \hat{\theta}$ 时, $\Pr(G | piv, \theta_i; \hat{\theta}) \leq \frac{1}{2}$; 在 $\theta_i > \hat{\theta}$ 时, $\Pr(G | piv, \theta_i; \hat{\theta}) \geq \frac{1}{2}$ 。

相关分析不免烦琐, 但是, 我们能很容易地推导出为使 $\hat{\theta} \in (0, 1)$ 存在, 这一博弈所需满足的条件。首先, $\Pr(G | piv, \theta_i; \hat{\theta}) \geq \frac{1}{2}$ 当且仅当:

$$\begin{aligned} & \frac{\pi f(\theta_i | G) F(\hat{\theta} | G)^{n-q-1} [1 - F(\hat{\theta} | G)]^{q-1}}{(1 - \pi) f(\theta_i | I) F(\hat{\theta} | I)^{n-q-1} [1 - F(\hat{\theta} | I)]^{q-1}} \\ &= \frac{f(\theta_i | G)}{f(\theta_i | I)} \frac{\pi F(\hat{\theta} | G)^{n-q-1} [1 - F(\hat{\theta} | G)]^{q-1}}{(1 - \pi) F(\hat{\theta} | I)^{n-q-1} [1 - F(\hat{\theta} | I)]^{q-1}} \\ &\geq 1 \end{aligned}$$

于是, 根据严格单调似然率条件, 如果 $\Pr(G | piv, \hat{\theta}; \hat{\theta}) = \frac{1}{2}$, 则 $\theta_i < \hat{\theta}$ 意味着

$\Pr(G | piv, \theta_i; \hat{\theta}) \leq \frac{1}{2}$, $\theta_i > \hat{\theta}$ 意味着 $\Pr(G | piv, \theta_i; \hat{\theta}) \geq \frac{1}{2}$ 。如果

$\Pr(G | piv, 0; 0) \leq \frac{1}{2} \leq \Pr(G | piv, 1; 1)$, 由于函数 $\Pr(G | piv, \cdot; \cdot)$ 连续, 所以

根据介值定理, 存在这样的折点。对很多这类博弈, 这些边界条件都得到满足。

在简单的两元信号模型中, 每个人使用相同投票规则且私人信息决定投票行为的均衡, 未必存在。但是, 在连续统模型中, 这种均衡一般都存在。如果采用一致通过的投票规则, 则模型的选择对投票结果有影响。在两元信号模型中, Feederson 和 Pesendorfer(1998)指出, 当参与者数量众多时, 一致通过的投票规则是唯一差的信息加总方式, 因为在均衡中, 每个陪审员在其他所有人都投票认定被告有罪的假设下投票。Meirowitz(2002)指出, 在连续统模型中, 一致通过的投票规则常常和其他的投票规则一样好。

6.5 应用: 公共品和不完全信息

在本节介绍的 Palfrey-Rosenthal 博弈中, 潜在的出资人并不了解其他参与者的出资成本。为使模型尽可能地接近第 4 章中的模型, 我们依然假定在至少有 k 个人出资时, 每个参与者得到 1 个单位的效用; 否则, 效用是 0。第 i 个参与者的出资成本为 c_i , c_i 在 $[0, 1]$ 上均匀分布。每个参与者了解自己的成本, 但不知道其他人的成本。

6.5.1 $k=1$

在我们所探讨的第一种情况中,公共品被生产出来的充要条件是有一人提供了资金。由于对所有的 i , $c_i \leq 1$, 所以存在 n 个贝叶斯 Nash 均衡, 在第 i 个均衡中, 只有参与者 i 提供了资金。和前面一样, 我们只讨论有相同成本的所有的参与者类型使用相同策略的对称均衡。就是说, 我们探讨折点策略上的均衡。在这种均衡中, 第 i 个参与者提供资金当且仅当 $c_i < \hat{c}_n$, 其中, $\hat{c}_n$ 是这个有 n 个参与者的博弈的均衡折点。^① 如果参与者 i 的所有对手都选择这一折点策略, 则他从提供资金中得到的效用为 $1 - c_i$ 。如果他没有提供资金, 则有两种情况: 如果有另外至少一个参与者提供了资金, 则 i 得到的效用为 1; 否则, 就为 0。由于 c 在 $[0, 1]$ 上均匀分布, 其他出资者以概率 $\hat{c}$ 提供资金, 所以没有其他任何人提供资金的概率为 $[1 - \hat{c}_n]^{n-1}$ 。于是, 第 i 个参与者从不提供资金的行为中得到的效用为 $1 - [1 - \hat{c}_n]^{n-1}$ 。因此, 只要:

$$[1 - \hat{c}_n]^{n-1} \geq c_i$$

参与者 i 就会提供资金。由于成本为 $\hat{c}_n$ 的参与者在自己的选择上肯定是无差异的, 所以, 在贝叶斯 Nash 均衡中:

$$[1 - \hat{c}_n]^{n-1} = \hat{c}_n$$

折点 $\hat{c}_n$ 在 n 上递减。如果并非如此, 则可以设 $\hat{c}_{n+1} \geq \hat{c}_n$ 。折点条件意味着 $[1 - \hat{c}_{n+1}]^n \geq [1 - \hat{c}_n]^{n-1}$ 。但是, 由于 $\hat{c}_n$ 和 $\hat{c}_{n+1}$ 都位于 0 和 1 之间, 所以, 这一不等式意味着 $\hat{c}_n > \hat{c}_{n+1}$ 是矛盾的。由于 $\hat{c}_n$ 在 n 上递减, 所以, 随着这个社会人数增加, 任何参与者提供资金的概率趋于 0。但是, 没有任何人提供资金的概率 $[1 - \hat{c}_n]^n$ 也趋于 0。所以在这一模型中, 虽然随着 n 上升, 任何一个参与者提供资金的概率收敛于 0, 但是, 提供资金的概率 $1 - [1 - \hat{c}_n]^n$ 收敛于 1。

6.5.2 $k > 1$

设公共品的生产需要多个人提供资金。我们依然假设参与者使用折点策略, 只在 $c_i \leq \hat{c}_n$ 时才提供资金。 x_{-i} 表示由 i 以外的所有人提供的资金数量。根据前一章的分析, 第 i 个人从提供资金的行为中得到的净效用是:

$$\Pr(x_{-i} = k - 1) - c_i$$

① 所有均衡都是折点策略上的均衡。给定 i 的任意策略, 设 p_{-i} 是至少有另一个参与者提供资金的概率。给定 p_{-i} , 第 i 个参与者从提供资金的行为中得到的期望效用是 $1 - c_i$, 从不提供资金的行为中得到的期望效用是 p_{-i} 。于是, 在 $c_i < 1 - p_{-i}$ 时, 第 i 个参与者的最优反应是提供资金; 在 $c_i > 1 - p_{-i}$ 时, 最优反应是不提供资金。这一最优反应便是折点策略。

由于每个人都以先验概率 $\hat{c}_n$ 提供资金, 所以:

$$\Pr(x_i = k-1) = \binom{n-1}{k-1} \hat{c}_n^{k-1} (1 - \hat{c}_n)^{n-k}$$

由于第 i 个参与者在自己的成本是 $c_i = \hat{c}_n$ 时, 在出资和不出资上是无差异的, 所以我们能得到 $\hat{c}_n$ 的隐性解:

$$\binom{n-1}{k-1} \hat{c}_n^{k-1} (1 - \hat{c}_n)^{n-k} = \hat{c}$$

这个解类似完全信息下的混合策略均衡解, 只是在这里, $\hat{c}_n$ 起着 σ^* 的作用。所以, 前一章的许多结论可以直接移植过来。

为简化符号, 设:

$$\Pi(\hat{c}_n) = \binom{n-1}{k-1} \hat{c}_n^{k-2} (1 - \hat{c}_n)^{n-k}$$

于是, 均衡条件为 $\Pi(\hat{c}_n) = 1$ 。对 $\Pi(\hat{c}_n)$ 求导, 得到:

$$\frac{\partial \Pi(\hat{c}_n)}{\partial \hat{c}_n} = \binom{n-1}{k-1} [(k-2) \hat{c}_n^{k-3} (1 - \hat{c}_n)^{n-k} - \hat{c}_n^{k-2} (n-k) (1 - \hat{c}_n)^{n-k-1}]$$

这个导数和 $-(\hat{c}_n n - k - 2\hat{c}_n + 2)$ 有相同的符号。就是说:

$$\text{如果 } \hat{c}_n \begin{cases} \geq \frac{k-2}{n-2} \\ < \frac{k-2}{n-2} \end{cases}, \text{ 则 } \Pi' \begin{cases} \geq \\ < \end{cases} 0$$

所以, 只要 $2 < k < n$, $\Pi(\hat{c}_n)$ 就递增至位于点 $\hat{c}_n = (k-2)/(n-2)$ 的唯一的最大值上, 然后在 $\hat{c}_n > (k-2)/(n-2)$ 的区间上, 逐渐递减。在 $\hat{c}_n \in \{0, 1\}$ 时, $\Pi(\hat{c}_n) = 0$ 。 $\Pi(\hat{c}_n)$ 的这些特点表明, 如果 $2 < k < n$ 且 $\max \Pi > 1$, 就存在两个均衡折点 $\hat{c}_H$ 和 $\hat{c}_L$ 。如果 $\max \Pi < 1$, 则在 $(0, 1)$ 中不存在折点上的均衡。^①

和前一章一样, 我们能很容易地证明, 随着 n 变得很大, 折点均衡消失。调整均衡条件, 得到:

$$\binom{n-1}{k-1} \hat{c}_n^{k-1} (1 - \hat{c}_n)^{n-k} = \hat{c}_n$$

对任意的 $\hat{c}_n$, 等号左边的部分:

$$p(n, \hat{c}_n) = \binom{n-1}{k-1} \hat{c}_n^{k-1} (1 - \hat{c}_n)^{n-k}$$

是在成功概率为 $\hat{c}_n$ 时, 在 $n-1$ 次试验中成功 $k-1$ 次的概率。这一概率函数具有

① 在 $\max \Pi = 1$ 的这种不大可能发生的事件中, 存在唯一的折点均衡。

的一个性质是,除非 $\hat{c}_n$ 收敛于预期成功率 $(k-1)/(n-1)$, 否则随着 n 趋向无穷大, $p(n, \hat{c}_n)$ 收敛于 0。^① 如果 $\hat{c}_n$ 不收敛于 $(k-1)/(n-1)$ 而 $p(n, \hat{c}_n)$ 收敛于 0, 则均衡条件要求 $\hat{c}_n$ 收敛于 0。如果 $\hat{c}_n$ 收敛于 $(k-1)/(n-1)$, 则它肯定收敛于 0, 因为 $(k-1)/(n-1)$ 收敛于 0。

作为习题,我们请读者证明, n 和 k 对 $\hat{c}_H$ 和 $\hat{c}_L$ 的影响,与前一章的对称混合策略均衡中 n 和 k 对 σ_L^* 和 σ_H^* 的影响基本相同。

6.6 应用: 候选人偏好上的不确定性

在前一章关于候选人竞选的 Hotelling 模型中,我们假定候选人是政策导向的,并且在中间选民的位置上存在着不确定性。现在,除了中间选民位置上的不确定性外,我们假设候选人对自己的政策偏好拥有私人信息。候选人 1 有理想点 $\theta_1 \in \{0, \frac{1}{2}\}$, 候选人 2 有理想点 $\theta_2 \in \{\frac{1}{2}, 1\}$ 。政策 x 带给候选人 i 的效用是 $u(x) = -(\theta_i - x)^2$ 。为简化分析,假设所有候选人类型的生成概率相等,且他们的类型相互独立。和前一章一样,中间选民的理想点是从 $[0, 1]$ 上的均匀分布中随机抽取的。候选人 1 的策略是映射 $s_1(\theta_1): \{0, \frac{1}{2}\} \rightarrow [0, \frac{1}{2}]$, 候选人 2 的策略是映射 $s_2(\theta_2): \{\frac{1}{2}, 1\} \rightarrow [\frac{1}{2}, 1]$ 。为简化分析,我们不考虑候选人从自己的类型区间外部选择策略的可能性。首先设候选人 2 使用的策略为 $s_2(\frac{1}{2}) = a$ 和 $s_2(1) = b$ 。对类型为 $\theta_1 = \frac{1}{2}$ 的候选人,任何位置 $s_1 < \frac{1}{2}$ 都劣于位置 $\frac{1}{2}$, 因为对给定的 a 和 b , 位置 $\frac{1}{2}$ 比位置 $s_1 < \frac{1}{2}$ 更有可能胜出。而且,对类型为 $\theta_1 = \frac{1}{2}$ 的候选人 1 来说,从赢得选举的角度看,位置 $\frac{1}{2}$ 比位置 $s_1 < \frac{1}{2}$ 更理想。因此, $s_1(\frac{1}{2}) = \frac{1}{2}$ 和 $s_2(\frac{1}{2}) = \frac{1}{2}$ 严格优于其他一切竞选纲领。

类型为 $\theta_1 = 0$ 的候选人 1 的最优反应,是以下最大化问题的解:

$$\max_{s_1} \left\{ -s_1^2 \left[\frac{s_1 + \frac{1}{2}}{4} + \frac{s_1 + b}{4} \right] - \frac{\left(\frac{1}{2}\right)^2}{2} \left[1 - \frac{s_1 + \frac{1}{2}}{2} \right] - \frac{b^2}{2} \left(1 - \frac{s_1 + b}{4} \right) \right\}$$

把目标函数对 s_1 求导,并让导数等于 0,得到一阶条件:

① 这一结论是大数法则的应用。根据大数法则,样本均值的极限分布只对随机变量的期望值赋以正的权重。

$$\frac{1}{8}b^2 - \frac{1}{2}bs_1 - \frac{1}{4}s_1 - \frac{3}{2}s_1^2 + \frac{1}{16} = 0$$

由于最优解应落于区间 $\left[0, \frac{1}{2}\right]$ 中, 所以这一解为:

$$s_1(0; b) = \frac{1}{12}\sqrt{4b+16b^2+7} - \frac{1}{6}b - \frac{1}{12}$$

候选人 2 的问题和候选人 1 的问题相同。就是说, 通过求解相关一阶条件, 我们能够找到 $s_2(1) = b$ 和 $s_1(0) = 1 - b$ 的均衡值。b 的均衡值满足下面的方程:

$$1 - b = \frac{1}{12}\sqrt{4b+16b^2+7} - \frac{1}{6}b - \frac{1}{12}$$

解这一方程, 得到 $b = 11/7 - \sqrt{106}/14 \simeq 0.836$ 。因此, 贝叶斯 Nash 均衡为 $s_1(0) = 0.164$, $s_1\left(\frac{1}{2}\right) = \frac{1}{2}$, $s_2\left(\frac{1}{2}\right) = \frac{1}{2}$, $s_2(1) = 0.836$ 。我们把类型为 $\theta_1 = 0$ 和 $\theta_2 = 1$ 的候选人的竞选纲领, 与完全信息博弈中候选人的理想点分别为 0 和 1 时的结果, 进行比较。你可能会认为, 不完全信息中两个候选人的竞选纲领会趋向一致, 因为每个候选人都认为自己可能与一位温和对手竞选。但这种认识并不全面。由于这两位候选人都是政策导向的, 所以他们宁愿输给温和对手而不愿输给极端对手。这一条件弱化了极端候选人走温和路线的激励。实际上, 在这个候选人的偏好不确定的博弈中, 候选人的竞选纲领甚至比其偏好为公开信息的博弈中的纲领更为分化。

6.7 应用: 竞选、竞赛和拍卖

到目前为止在我们所介绍的竞选模型中, 候选人只能在某个连续统中选择自己的政策。这种限制忽略了现实竞选活动中的许多策略性选择。本节探讨的模型建立在关于竞赛的经济学模型基础上。在竞赛模型中, 候选人选择自己的努力水平, 这种努力是要发生成本的。候选人投入的努力程度越高, 胜出的可能性就越大。这里探讨货币在竞选中的作用。假设有 n 个人竞选公职, 参与者集合为 $N = \{1, \dots, n\}$ 。候选人竞相筹措资金并用筹得的资金购买广告。设 $a_i \in \mathbb{R}_+$ 是第 i 个候选人的广告量。给定 $a = (a_1, \dots, a_n)$, 选举结果由函数 $p(a): \mathbb{R}_+^n \rightarrow N$ 决定, 这里的 p 是弱递增函数。设 $p(a) = \arg \max_{i \in N} a_i$, 即这一公职给了广告量最大的候选人。^① 候选人 i 的效用决定于谁是胜出者、自己的广告量 a_i 和自己对这一公职的价值认定 $\theta_i \in [0, 1]$ 。每个候选人对这一公职的价值认定是私人信息, 是从 $[0, 1]$ 上的均匀分布中独立抽取的。具体地说, 第 i 个候选人的效用为 $u_i(a) = \theta_i 1_{\{p(a)=i\}} - a_i$, 其中, $1_{\{p(a)=i\}}$ 是指示函数, 如果 $p(a) = i$, 它

① 在这一模型中, 还可以把 a_i 理解为候选人竞选公职过程中投入的精力或时间。

的值是1, 否则是0。所有候选人同时选择自己的广告支出水平 a_i , 各自的收益随即得以实现。^①贝叶斯 Nash 均衡是个函数, 对每个候选人, 它把 $\theta_i \in [0, 1]$ 映射到 $a_i \in \mathbb{R}_+^1$ 。这里只探讨有相同类型的两个候选人选择相同的广告支出水平的对称均衡。

直接求解连续的策略函数常常很难。所以, 我们玩点儿技巧。我们假设策略函数有具体的函数形式。我们对任意的自由参数求解, 证明求得的解构成均衡。在这里, 我们猜测参与者 $j \neq i$ 使用的策略为 $a_j(\theta_j) = b\theta_j^c$, 其中的 b 和 c 是待定参数。如果参与者 $2, \dots, n$ 都使用我们所猜测的这种策略, 并且候选人 1 选择 a_1 , 则他胜出的概率为 $\Pr\{\max_{j \neq i} b\theta_j^c < a_1\}$ 。这一概率为

$$\Pr\left\{\max_{j \neq i} \theta_j < \left(\frac{a_1}{b}\right)^{\frac{1}{c}}\right\} = \left(\frac{a_1}{b}\right)^{\frac{n-1}{c}}$$

类型为 θ_1 的参与者 1 从行动 a_1 中得到的期望效用于是为:

$$\left(\frac{a_1}{b}\right)^{\frac{n-1}{c}} \theta_1 - a_1$$

对 a_1 求导, 得到一阶条件:

$$\theta_1 \frac{n-1}{cb} \left(\frac{a_1}{b}\right)^{\frac{n-1}{c}-1} = 1$$

求得 a_1 为:

$$a_1 = b \left(\frac{cb}{(n-1)\theta_1} \right)^{\frac{c}{n-1-c}} \quad (6.3)$$

在这里, 我们首先猜测参与者 $j = 2, \dots, n$ 都使用形式为 $a_j(\theta_j) = b\theta_j^c$ 的策略, 进而发现参与者 1 的最优反应也是使用这一形式的策略。于是, 通过求解 b 和 c 的值, 我们就能找到均衡。 b 和 c 满足:

$$b\theta_1^c = b \left(\frac{cb}{(n-1)\theta_1} \right)^{\frac{c}{n-1-c}}$$

于是得到 $b = \frac{(n-1)}{n}$ 和 $c = n$ 。将这两个结果代入方程(6.3)等号左边并简化, 可得到:

$$a_1 = \frac{n-1}{n} \theta_1^n$$

这一结果表明, $b = \frac{(n-1)}{n}$ 和 $c = n$ 与均衡相一致。因此, 在对称贝叶斯 Nash 均衡中, 每个参与者选择:

① 这一模型等价于类型独立的最高价格全支付拍卖。

$$a_i(\theta_i) = \frac{n-1}{n} \theta_i^n$$

从这一均衡中,我们得到一些结论。首先,在均衡中,候选人的广告支出水平与他在这个职务的价值认定正相关。其次,广告支出水平和这个职位的价值之间的关系,决定于候选人的数量。对上式求导,得到:

$$\frac{\partial a_i(\theta_i)}{\partial \theta_i} = \theta_i^{n-1}(n-1)$$

$$\frac{\partial^2 a_i(\theta_i)}{\partial \theta_i \partial n} = \theta_i^{n-1}(n \ln \theta_i - \ln \theta_i + 1)$$

其中,交叉导数是负的(因为对 $\theta_i \in (0, 1)$, $\ln \theta_i < 0$)。就是说,随着候选人数量上升,均衡策略变得更加平坦。并且,随着 n 增大,候选人广告支出的上界收敛于 0 (就是说,随着 n 增大, $a_i(1)$ 趋向 0)。这一结论表明,在有許多人参与竞选时,候选人有最高的 θ 值的概率收敛于 0;在不可能胜出的竞选中,候选人不愿意投入太多资源。

不难看出这一模型和拍卖模型之间的关系。在这一博弈中,无论是否胜出,候选人都要发生负效用 a_i 。我们还可以做一些发展,例如所有候选人都宣布,如果自己当选,则兑现竞选时承诺的支付;或者,利益集团向委员会主席承诺,如果自己所偏好的提名者得到确认,就提供捐款;在后面关于机制设计的一章中,我们将探讨许多种此类拍卖模型。

6.8 贝叶斯 Nash 均衡的存在性

贝叶斯 Nash 均衡肯定存在吗? 贝叶斯标准式博弈只是标准式博弈的一种特殊情况:所有的参与者类型同时选择策略,收益被定义为参与者在策略组合上的期望效用。所以,这种均衡的存在性所要求的条件,和 Nash 均衡的存在性要求的条件相似。对类型空间和行动空间为有限集的博弈,我们可以用前面的结论来证明混合策略上贝叶斯 Nash 均衡的存在。

给定贝叶斯博弈 $\langle N, S, \Theta, u, p \rangle$, 其中 N 、 S 和 Θ 都是有限集。不失一般性,我们把类型集表示为 $\Theta_i = \{\theta_i^1, \dots, \theta_i^n\}$, 并定义标准式博弈 Γ' : 设 $N' = \{\theta_i^1, \dots, \theta_i^n, \theta_i^1, \dots, \theta_i^n\}$ 。在这个标准式博弈中,所有有下标 i 的参与者的策略空间为 $S'_i = S_i$ 。 $\Theta_{-i} = \times_{j=N \setminus \{i\}} \Theta_j$ 表示原始贝叶斯博弈中参与者 $N \setminus \{i\}$ 的可能的类型组合集。给定博弈 Γ' 中的策略组合 $s^+ = (s_1^1, \dots, s_i^1, \dots, s_i^n, \dots, s_n^1, \dots, s_n^n) \in \times_{i=1}^n S'_i$, 令 $s_i^+(\theta_i = \theta_i^j) = s_i^j$, 我们可以将这一策略定义为博弈 Γ 中的策略。于是,利用原始贝叶斯博弈中的期望效用概念,我们就能定义参与者 θ_i^j 的效用为:

$$v_i^j(s^+) = EU_i(s_i^+(\theta_i), s_{-i}^+(\cdot); \theta_i^j)$$

有了这个有限标准式博弈 $\Gamma' = \langle N', S', v \rangle$, 就有下面的关于存在性的结论了。

命题 6.1 给定贝叶斯博弈 $\langle N, S, \Theta, u, p \rangle$, 其中 N 、 S 和 Θ 是有限集, 则存在混合策略上的贝叶斯 Nash 均衡。

证明: 根据 Nash 定理(定理 5.4), 有限博弈 $\langle N', S', v \rangle$ 有混合策略 Nash 均衡。用 σ_i' 表示在这一混合策略 Nash 均衡中策略集 S_i 上的概率分布。由于策略组合 σ 满足 Nash 均衡所要求的条件, 所以, 在其中以概率 $\sigma_i'(s')$ 使得 $\phi_i(\theta_i) = s'$ 的策略, 满足条件 6.1。□

6.9 习题

习题 6.1 在陪审团投票博弈, 设 $p = \frac{3}{4}$, $\pi = \frac{2}{3}$ 。分析这一博弈的贝叶斯 Nash 均衡集的特征。如果陪审团采用的不是多数通过而是一致通过的投票规则, 就是说, 只有在所有人都认定被告有罪时, 被告被判决有罪; 如果有一人投票认定其无罪, 则被告被宣判无罪。分析这一博弈中的贝叶斯 Nash 均衡 (依然假设 $p = \frac{3}{4}$, $\pi = \frac{2}{3}$)。

习题 6.2 在陪审团投票模型中, 设类型集为连续统。证明方程(6.2)。

习题 6.3 在 Palfrey-Rosenthal 模型中, 假设要生产出公共品, 需要 k 个人提供资金。如果出资者数量小于 k , 所提供的资金全额退回。这一调整如何影响折点 $\hat{c}$ 的值? 如果出资者的数量超过 k , 多余的资金被随机地退回给参与者, 这时的结果又会怎样呢?

习题 6.4 在 Palfrey-Rosenthal 模型中, 假设要生产出这件公共品, 需要 k 个人提供资金。就 n 和 k 的变化对内生值 $\hat{c}_n$ 和 $\hat{c}_k$ 的影响, 进行比较静态分析。

习题 6.5 在竞选博弈中, 设在候选人的偏好上存在着私人信息, 同时假设每个人都知道中间选民的位置为 $\frac{3}{4}$ 。分析贝叶斯 Nash 均衡的特征。

7

扩展式博弈

在标准式博弈中,所有参与者同时选择自己的策略,所以标准式博弈是静态博弈。但是,在政治科学的许多应用中,参与者按时间顺序选择自己的策略。我们可以把这类情况表述为标准式博弈,但使用扩展式博弈常常更容易,更能得到满意的效果,因为扩展式博弈明确引入了时间因素。

为方便理解扩展式博弈,我们看一个例子。A 国是 B 国统治的殖民地。B 国通过控制 A 国的油田和对 A 国居民直接征税而获得收入。

在第一个阶段,A 国决定是“革命,推翻 B 国统治”(R)还是“接受现状”(C)。如果 A 国决定推翻 B 国殖民,则在第二个阶段,B 国决定是“允许 A 国独立”(C)还是“镇压 A 国革命”(S)。如果 B 国镇压 A 国的革命,局势很快恶化为战争。一旦发生战争,A 国获胜概率是 p 。双方争夺的是利润滚滚的油田,它给控制者带来 4 个单位的收益。

在第二个阶段,如果 B 国不采取镇压措施,则发起革命给 A 国造成 1 个单位的成本;如果 B 国镇压革命,则双方各付出 6 个单位的成本。如果 A 国接受现状,B 国能继续向 A 国居民征税(T),获得 2 个单位的收入;它也可以取消税收(E)。表 7.1 给出了各种可能的结果所产生的收益,其中的第一个数字为 A 的收益。

表 7.1 革命博弈

A 不改变现状,B 取消税收,(0, 4)
A 不改变现状,B 继续征税,(-2, 6)
A 革命,B 允许 A 独立,(3, 0)
A 革命,B 镇压革命, $(4p - 6 - 2(1 - p), 6(1 - p) - 6)$

如果把这一博弈表示为标准形式,我们就会忽视 B 在做出自己的决策时已知道 A 的选择的事实。表述这一博弈的更好的方式是用图 7.1 中的博弈树。博弈树由表示所有已做出的决策的节点构成。换句话说,节点为博弈中已发生的历史。在博弈开始时,存在初始节点。在每个节点上,有树枝表示在此节点做出选择的参与者可以选择的行动。每个树枝指向下个阶段的节点。博弈结束是用终节点表示,终节点给出了博弈中每个参与者的收益。

A 国在初始节点做出选择。从这个节点出发,有两个树枝,分别对应于 R 和 C。在 A 国的革命决策之后,B 在 C 和 S 之间做出选择。在 A 选择 C 之后,B 在 T 和 E 之间做出选择。在每个终节点旁边,给出了相对应的收益。

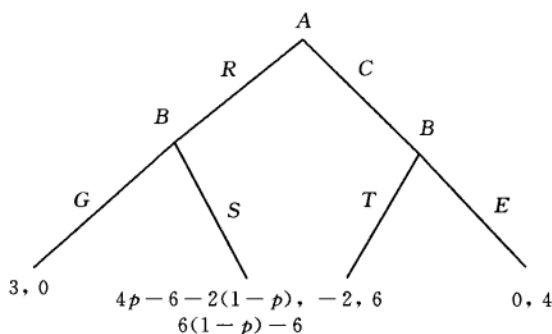


图 7.1 革命博弈

就像用矩阵表示标准式博弈一样,博弈树表示扩展式博弈。扩展式博弈的构成因素如下:

(1) 参与者集合 N 。

(2) 历史集 H 。 H 中的元素对应着博弈树的节点。 H^T 是由所有的终点历史构成的集合。习惯上,初始节点被表示为 $H^0 = \emptyset$, 即空集。

(3) 映射 $p(h): H \setminus H^T \rightarrow N$ 赋予每个非终点历史 h 一个参与者,由他在节点 h 做出决策。

(4) 对每个 h ,有一个行动集 $A(h)$ 。在历史 h 之后, $p(h)$ 必须从这个集合中选择行动。这种行动集可能包含行动上的随机选择。

(5) 构成历史集的分割的信息集 $I \subseteq H \setminus H^T$ 。如果 $h \in I$ 且 $h' \in I$, 则 $p(h)$ 不清楚自己是位于节点 h 还是节点 h' 。 h 和 h' 如果位于同一个信息集中, 则 $p(h) = p(h')$ 。在前面的博弈中, 每个参与者在轮到自己行动时, 都知道所发生的历史, 因此每个信息集都包含唯一一个元素。在所有信息集都是单元素集合时, 这一博弈有完全且完美信息(或者简单地说, 有完美信息)。在后面我们将放松这一假设, 允许参与者没有观察到其行动之前的所有行动。就是说, 一些信息集可能包含多个元素。这就是完全但非完美信息博弈(或者简单地说, 非完美信息博弈)。我们要求历史满足某些条件以保证它们构成良好形态的博弈树。^①

① 具体地说, 节点集上有弱序“……在……之前”。节点 a 如果在节点 b 之前, 我们就将两者之间的关系表示为 $a \rightarrow b$ 。我们说过, 弱序具有传递性, 所以在我们的扩展式博弈中, 没有循环。另外, 为避免多个节点指向同一节点, 我们假定, 如果 $h \rightarrow h'$ 且 $h'' \rightarrow h'$, 则或者 $h \rightarrow h'$ 或者 $h'' \rightarrow h$ 。存在唯一一个节点(初始节点) H^0 , 对每个 $h \in H$, $H^0 \rightarrow h$ 。在任何终点历史 $h \in H^T$ 上, 不存在任何历史 h' 使得 $h \rightarrow h'$ 。

(6) 收益 U 是一组贝努里效用函数。每个 $i \in N$ 有效用函数 $u_i(h): H^T \rightarrow \mathbb{R}^1$ 。

我们可以把有限扩展式博弈 Γ^E 定义为 $(N, H, p(\cdot), U)$ 。在扩展式博弈中, 一个策略便是一个完整的行动方案。它为每个参与者规定了在每个由此参与者采取行动的历史上, 他可以采取的行动。策略的正式定义如下:

定义 7.1 对扩展式博弈 Γ^E , 参与者 $i \in N$ 的策略是个映射 $s_i(h): H_i \rightarrow A(h)$ 。其中, 如果 h 和 h' 位于同一个信息集中, 则 $s_i(h) = s_i(h')$; H_i 是由使 $p(h) = i$ 的历史 $h \in H$ 构成的集合。策略组合则是映射 $s(h): H \setminus H^T \rightarrow A(h)$ 。如果 $h \in H_i$, 则 $S(h) = S(h_i)$ 。

利用这一定义, 我们写出在革命博弈中, 两个参与者各自的策略集。由于 A 只在初始节点上采取行动, 他的策略集为 $\{R, C\}$ 。对 B 来说, 她的每个策略必须明确她在每个信息集上的行动。她的策略集是 $\{G \text{ 和 } T, G \text{ 和 } E, S \text{ 和 } T, S \text{ 和 } E\}$, 其中, “G 和 T”说的是在 R 之后允许 A 国独立, 在 C 之后对 A 国征税 T。

一旦明确了策略, 就能够把这一博弈表述为标准式博弈, 见表 7.2。

表 7.2 标准式革命博弈

$B \backslash A$	R	C
G 和 T	0, 3	6, -2
G 和 E	0, 3	4, 0
S 和 T	$4p - 6 - 2(1 - p), 6(1 - p) - 6$	6, -2
S 和 E	$4p - 6 - 2(1 - p), 6(1 - p) - 6$	4, 0

利用这一形式能很容易地证明, 这个博弈有三个 Nash 均衡, 前两个均衡分别是策略组合 $(R, G \text{ 和 } T)$ 与 $(R, G \text{ 和 } E)$ 。在这两个均衡中, A 发动革命, B 允许 A 独立。第三个 Nash 均衡是策略组合 $(C, S \text{ 和 } T)$ 。在这个均衡中, B 镇压革命的威胁阻止了 A 发动革命。

从这个例子可以看出在动态博弈中, Nash 均衡概念的一些局限性。具体地说, 第二个和第三个均衡不大可能出现。在第二个均衡中, 在 A 做出接受 B 的统治的决策后, B 取消税收。但是在这一信息集上, B 没有动机取消税收。换句话说, Nash 均衡允许在“均衡路径以外”的历史上, 有非理性的行为。就第二个均衡而言, 这一问题似乎无足轻重, 因为它所规定的行为与第一个均衡中的行为完全一样。但是, 在第一个均衡中, 所有的信息集上的行为都是理性的。再看第三个均衡。假设 A 选择革命而偏离这一均衡。在他的偏离之后, 对任意的 p , B 应该选择 G。这就是说, B 镇压革命的威胁是不可信的, 因为对他来说, 这样做是不合理的。这一和平结果是由不满足序贯理性的行为在支持, 因为它没有在每一个信息集上最大化此参与者的效用。

在下面两节中, 我们讨论适用于动态博弈的 Nash 均衡的提炼, 这种提炼消除了不满足序贯理性的策略。

7.1 反向归纳法

求解完美信息动态博弈的最常用方法是反向归纳法。在这种方法中,在最后每个节点上采取行动的参与者选择最大化其效用的行动。此前一个参与者,知道最后一位采取行动的参与者在每个节点上选择最优行动;在此前提下,他最优地选择自己的行动。这一过程不断地持续下去,直至每个参与者,在其后所有参与者都做出最优选择的前提下,做出自己的最优选择。

在革命博弈中,我们用反向归纳法求解。首先,我们要求 B 在自己的每个节点做出最优选择。在节点 **R**, B 从“允许 A 独立”的决策中得到更高效用。在节点 **C**,“征税”的决策能给 B 带来最大效用。A 预期 B 做出理性决策,于是选择“革命”。因此,用反向归纳法得到的解是(**R**, **G** 和 **T**),它是不包含不满足序贯的理性行为的 Nash 均衡。

这一方法虽然易于使用,但是它的正式表述却很烦琐。我们看几个例子。

7.1.1 例子:蜈蚣博弈

图 7.2 中的博弈为蜈蚣博弈,它在实验经济学中被广泛使用。两个参与者轮换着在“下”(D)和“左”(L)之间做出选择。如果选择 **D**,则博弈结束;如果选择 **L**,则博弈继续进行直至第五个阶段。实验人员觉得这一博弈有用的一个原因是,参与者 1 可能会不断地选择 **L**,希望借此到达第五个阶段而得到最大收益 6。但是,这种策略不是序贯理性的,在反向归纳过程中,无法存活下来。我们从第 5 个阶段开始由后向前分析。在第五个阶段,参与者 1 选择 **D**。回到第四个阶段,参与者 2 知道参与者 1 在最后阶段会选择 **D**,从而使参与者 2 只能得到 4 个单位的收益,所以,他会选择 **D**,借此实现 5 个单位的收益。回到第三个阶段,参与者 1 意识到,**L** 产生 2 个单位的收益而 **D** 产生 3 个单位的收益。所以,他选择 **D**。持续这一过程直至回到第一个阶段,我们看到,参与者 1 理性地选择 **D**。事实上,在反向归纳法中存活下来的唯一的策略组合是 (**D**, **D**, **D**, **D**, **D**)。

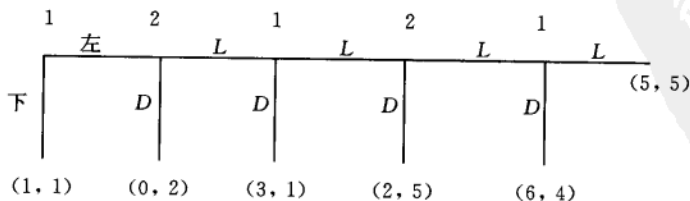


图 7.2 蜈蚣博弈

7.1.2 例子:序贯谈判

讨价还价模型在政治博弈论中的应用愈益重要。我们将在后面用一章内容探讨这类模型。这里只讨论一个最简单的模型。有两个参与者1和2,他们就1元钱在两人之间的分配而讨价还价。在第一个阶段,参与者1提出一个分配方案:他自己得到 x_1 而参与者2得到剩下部分,即 $x_2 = 1 - x_1$ 。如果参与者2接受这一方案,博弈结束。如果参与者2不接受这一方案,这1元钱所具有的价值就下降到 δ , $1 > \delta > 0$ 。施加这一假设是为了反映参与者缺乏耐性——他们希望尽早达成一致。在第二个阶段,参与者2提出方案:他自己得到 x_2 而参与者1得到剩下的部分,即 $x_1 = \delta - x_2$ 。如果参与者1接受这一方案,博弈结束。如果参与者1不接受这一方案,开始时的1元钱的值下降到零,两个参与者的所得为0。为简化分析,设每个参与者的收益函数为 $u_i(x_i) = x_i$ 。

在这一博弈中,存在许多Nash均衡。实际上,任何一种分配方案都能够得到Nash均衡策略的支持。我们看下面的策略组合:

参与者1:提议 $x_2 = z$ 。这一方案如果被拒绝,在第二阶段里,他就拒绝参与者2提出的一切方案。

参与者2:如果 $x_2 \geq z$,就在第一阶段里接受参与者1的方案;否则就拒绝这一方案,然后在第二阶段里提出方案 $x_2 = \delta$ 。

显然,给定 $z \leq 1$,参与者1的最优反应是在第一个阶段提议 $x_2 = z$ 。否则,参与者1得到的收益为0。参与者2的最优反应是接受 z 。但是,上述策略不满足序贯理性。对参与者1来说,拒绝第二阶段里的一切方案的做法不是理性行为。他应该接受给他带来的收益至少为0——0是拒绝方案而带给他的收益——的任何方案。这就意味着在第二阶段中,参与者2可以提议给自己几乎所有的 δ 。所以在第一阶段中,参与者1提供给参与者2的收益至少要和参与者2从自己在第二个时期中提议的分配方案中得到的收益一般多。下面的策略与反向归纳法相一致:

参与者1:提议 $x_2 = \delta$ 。如果这一方案被拒绝,则在第二阶段里,接受一切方案。

参与者2:如果 $x_2 \geq \delta$,就在第一阶段里接受参与者1提议的方案;否则就拒绝。在第二阶段里提出方案 $x_2 = \delta$ 。

在阶段2中,对参与者1来说,得到收益0并不比拒绝这一方案得到的收益低。因此,参与者2应提出的最优方案是 $x_2 = \delta$ 。回到第一阶段,参与者1提供给参与者2的收益至少必须为 δ ,只有这样才能防止后者拒绝方案。因此,在第一阶段中,均衡方案是 $x_1 = 1 - \delta$, $x_2 = \delta$ 。

7.2 完全非完美信息动态博弈

在目前为止所分析的模型中,在每个历史上行动的参与者,都知道此前所发生

的一切行动,并因此能推断出自己是在哪个节点上采取行动。换句话说,每个信息集中只包含一个元素。在下面的模型中,信息集中包含多个历史。这种形式的博弈有非完美信息。在某些行动或者没有被观察到或者同时进行时,这种情况就会出现。

我们分析行政官员 B 和政治家 P 之间的一个简单博弈。在这个博弈中,一些行动无法被观察到。政府官员从 $\{H, L\}$ 中选择执法力度, H 和 L 分别表示执法力度高和低。在执法力度高时, B 就得承担执法成本 $c > 0$; 在执法力度低时,执法成本为 0。为使分析简单,设 B 从执法活动中得不到任何效用。于是,在他选择 H 时,他得到收益 $-c$; 在他选择 L 时,他得到收益 0。在 H 和 L 之间,政治家偏好前者。他的效用函数为: $u_P(H) = 1$, $u_P(L) = 0$ 。P 除非进行监督,否则无法观察到 B 的执法力度;监督成本为 $1 > k > 0$ 。如果 B 被发现执法松弛,就会被罚款 f 并被强制选择 H 。

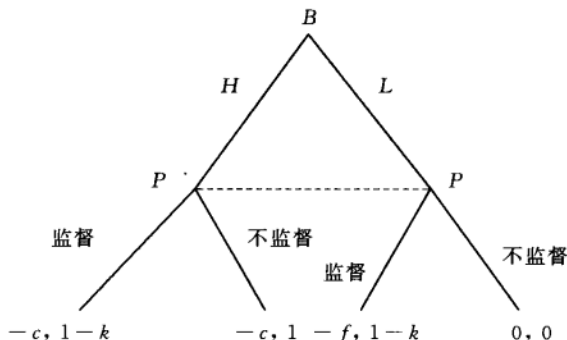


图 7.3 执法博弈

在决定是否监督时, P 不知道历史为 H 还是为 L 。在图 7.3 的扩展式博弈中, 连接 H 和 L 的虚线表示它们位于同一个信息集中。

由于 P 没有观察到 B 的行动,所以在每个节点上,它只能采取相同的行动。而 B 的一个策略便是在初始节点做出的一个选择,所以,这一博弈可以表示为表 7.3 中的标准形式。

在这一博弈中,不存在纯策略 Nash 均衡。如果 B 选择 H , P 的最优反应是不监督;但是,如果 P 选择不监督, B 的最优反应是 L ;在 B 选择 L 时, P 愿意进行监督。在这一博

表 7.3 监督博弈

B \ P	监 督	不监督
H	$-c, 1-k$	$-c, 1$
L	$-f-c, 1-k$	$0, 0$

弈的混合策略均衡中, B 以概率 $1-k$ 选择严格执法, P 以概率 $\frac{c}{f}$ 选择监督。

这个例子表明,在包含顺序行动的博弈中,如果存在没有被第二个参与者观察到的行动,这种博弈就和同时行动博弈完全相同。我们再用一个熟悉的例子说明如何用扩展式博弈表示同时行动博弈。在囚徒两难博弈中,两个嫌犯得决定是否

认罪。我们可以把这个博弈表示为扩展形式,让嫌犯1首先行动,然后如图7.4中那样,把“认罪”和“不认罪”放到参与者2的同一个信息集中。

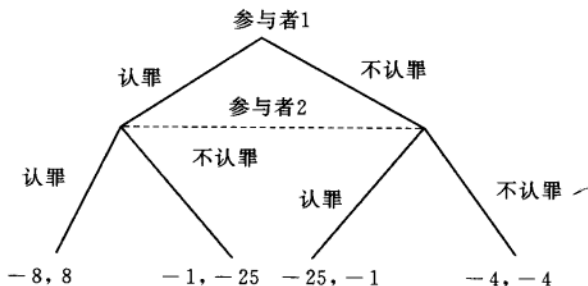
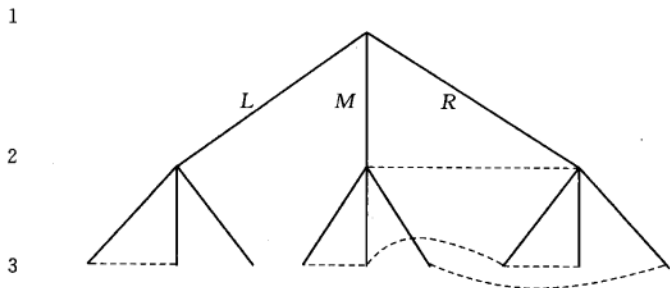


图 7.4 扩展形式的囚徒两难博弈

我们用图7.5中的三阶段博弈来说明扩展形式所具有的灵活性。每个参与者有三个行动可供选择:左(L)、中(M)和右(R)。如果参与者1选择L,这一行动被所有参与者观察到。如果参与者1选择R或M,参与者2就观察不到参与者1的行动。参与者2因此有两个信息集{L}和{M, R}。参与者3也无法完全观察到参与者1和参与者2的行动,因此他有四个信息集:{LL, LM}、{LR}、{ML, MM, RL, RM}和{MR, RR}。



在第一个阶段,参与者1只有一个单元素信息集。

在第二个阶段,参与者2有两个信息集。

在第三个阶段,参与者3有四个信息集。

图 7.5 复杂的信息集

非完美信息博弈的困难之处是,我们无法应用反向归纳法,因为参与者并不始终知道博弈到达了哪一个历史。对此问题有一种解决办法。首先要看到,扩展式博弈的某些部分可以被看作是独立的博弈。序贯理性概念要求所有参与者在每个这种小博弈中,使用Nash均衡策略。博弈内部的这种小博弈,被称为子博弈。子博弈满足下列标准:

(1) 它开始于某个单元素信息集。换句话说,在子博弈的初始节点或者说第一个节点上,采取行动的参与者知道自己在哪一个节点上。

(2) 它包含此初始节点之后的所有节点而不包含其他。

(3) 它没有切割任何信息集。历史 h 和 h' 如果在同一个信息集中, 它们就在同一个子博弈中。

图 7.5 中的博弈有三个子博弈: 原始博弈、从 L 开始的博弈和从历史 LR 开始的博弈。构成所有子博弈中的 Nash 均衡行为的策略组合, 被称为子博弈精炼均衡, 有时缩写为 SPNE。在定义 SPNE 时, 受限制的策略组合的概念能使定义变得简单。给定策略组合 $s(\cdot)$ 和历史集为 $\dot{H}'$ 的子博弈, $s(\cdot)$ 在这个子博弈上的限制, 是映射 s' ——它的定义域是 H' , 且对每个 $h \in H'$, $s'(h) = s(h)$ 。

定义 7.2 给定扩展式博弈 Γ^E , 策略组合 $s(\cdot)$ 是子博弈精炼 Nash 均衡 (SPNE), 如果在 Γ^E 的每个子博弈中, 受限制的策略组合 $s(\cdot)$ 是这个子博弈的 Nash 均衡。

下面的定理指出, 在有限博弈中存在 SPNE。

定理 7.1 每个有限扩展式博弈都有 SPNE。而且, 如果没有参与者在任意两个终点历史之间无差异, 则这个 SPNE 唯一。

我们来分析一个投票问题。设有三个人参与投票, $N = \{1, 2, 3\}$ 。他们依照以下程序对三个政策方案 x 、 y 和 z 进行投票: 首先在多数通过的规则下在 x 和 y 之间做出选择; 然后在多数通过的规则下, 就前一轮投票中胜出的方案与 z 进行投票。最后胜出的方案将被作为政策实施。在每个投票阶段中, 三个人同时投票。图 7.6 描述了这一博弈。

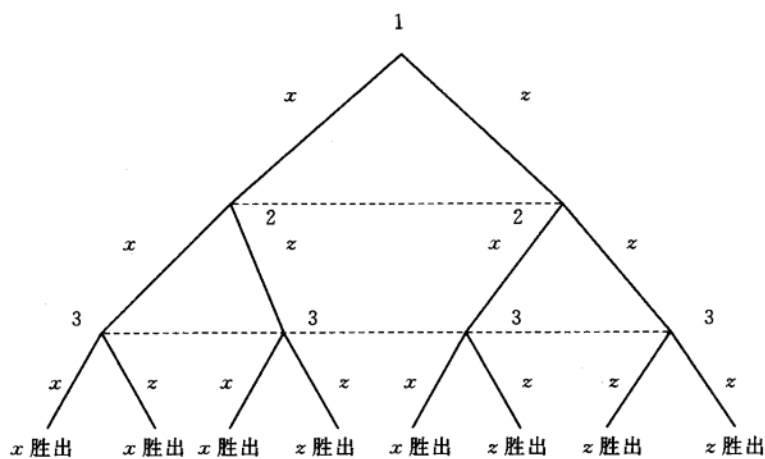
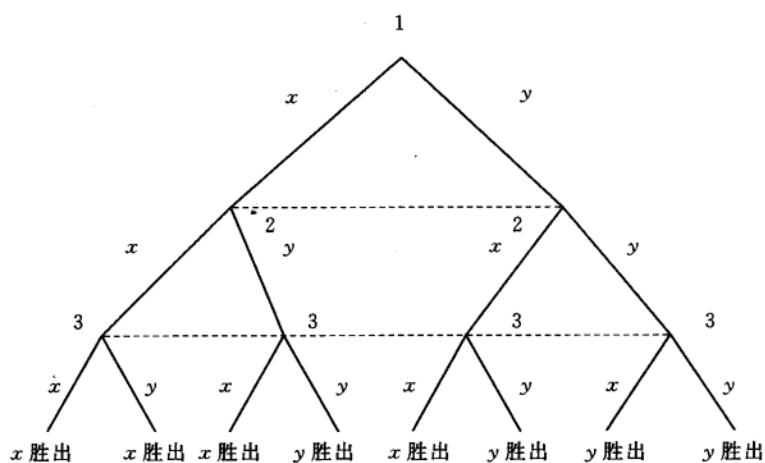
参与者在三个政策上的偏好分别为 xP_1yP_1z 、 yP_2zP_2x 和 zP_3xP_3y 。如果我们应用子博弈精炼标准并要求不使用弱劣策略, 则在最后一个阶段, 参与者 2 和参与者 3 投票给 z 而不是 x 。还有一种可能是, 参与者 1 和参与者 2 投票给 y 而不是 z 。

于是, 在第一个阶段当参与者就 x 和 y 投票时, 他们知道自己实际上是在复杂的等价方案 z 和 y 之间做出选择。^①于是, 参与者 1 和参与者 2 投票给 y 而不是 x 。参与者 1 虽然在 x 和 y 之间偏好 x , 但是在 y 和 x 之间, 他仍会策略性地投票给 y 。因为他知道, 投票给 x 实际上等同于投票给 z , 而这不是他希望出现的结果。

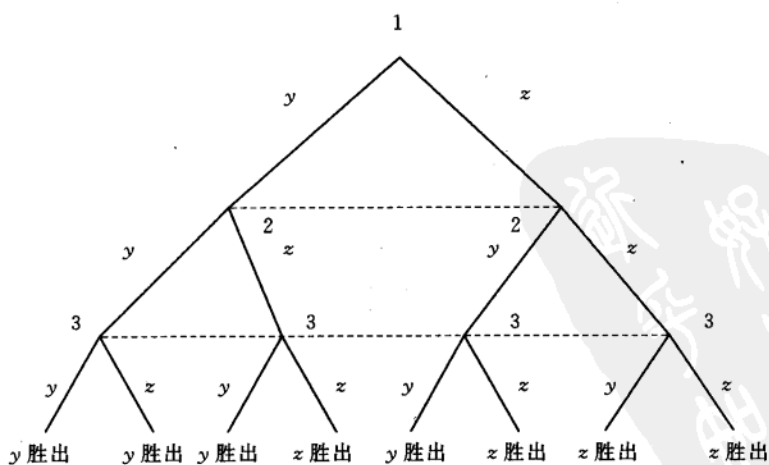
如果允许参与者使用弱劣策略, SPNE 集可能会很大。任何一致的投票行为, 都是第二个阶段中每个子博弈的 Nash 均衡。因此, 找出第二阶段的每个子博弈的 Nash 均衡策略, 就可以构造出许多 SPNE。

我们再用另一个模型说明 SPNE。该模型类似 Weingast(1997)解释法治社会演变的模型。在这个博弈中, 有一个统治者 R , 他决定是否对两个社会团体 A 和 B 中的一个剥削财富 x 。在观察到哪个团体被剥削后, A 和 B 同时决定是否挑战这位统治者。每个团体都要发生挑战成本 c 。如果两个团体都挑战, 则统治者的剥

① Richard McKelvey 和 Richard Niemi(1978)在对策略性投票的研究中, 提出了“复杂的等价”这一术语。



x 胜出的子博弈



y 胜出的子博弈

图 7.6 序贯投票博弈

削企图就会落空,每个团体得到收益 b 。如果挑战成功,还会迫使统治者付出成本 k 。如果没有团体挑战或者只有一个团体挑战,则统治者的剥削企图就会成功。图 7.7 给出了这一博弈的扩展形式。

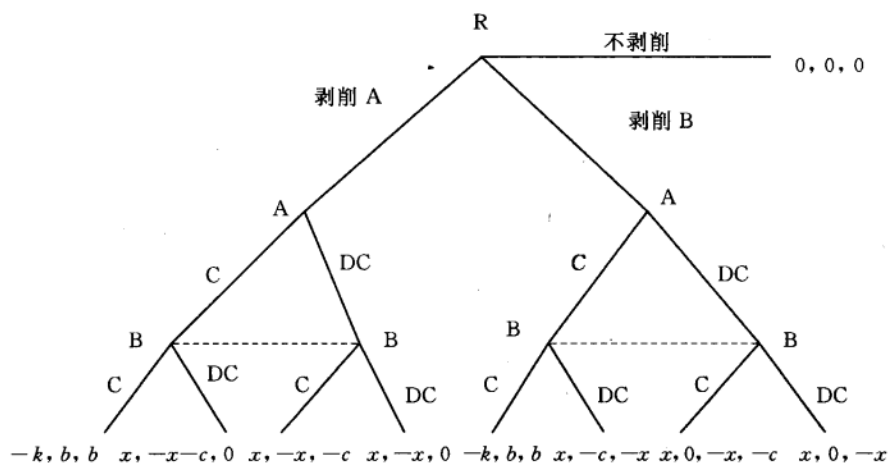


图 7.7 法治博弈

首先对自统治者“剥削 A”的决策开始的子博弈,计算 Nash 均衡。从 A 和 B 的角度看,这一子博弈的标准形式为表 7.4 所示。

这个子博弈中有两个纯策略 Nash 均衡:在一个均衡中,两个团体都挑战统治者;在另一个均衡中,没有人挑战统治者。它还有一个混合策略均衡,但为简化分析,我们不考虑这一均衡。由于自统治者“剥削 B”的企图开始的子博弈与上述子博弈对称,所以它也有与之相同的两个纯策略均衡。

表 7.4 剥削博弈

B \ A	挑 战	不挑战
挑 战	b, b	$-c, -x$
不挑战	$0, -x-c$	$0, -x$

再来分析这个博弈的第一阶段。R 意识到,在第二个阶段的每个子博弈中,所进行的必然是某个 Nash 均衡。如果他认为在第二个阶段的两个子博弈中两个团体都会挑战他,他的最优反应是不剥削。如果在子博弈中有一个团体挑战他而另一个团体没有挑战他,他就会剥削没有挑战它的那个团体。如果在两个子博弈中都没有团体挑战它,统治者就会任选一个团体进行剥削。所以,存在五个纯策略 SPNE。Weingast 认为,建立法制社会的关键是, A 和 B 协调到某个 Nash 均衡上,在这个均衡中,两个团体挑战统治者的剥削企图。

通过反向归纳求得的解只是 SPNE 的特例。在完美信息博弈中,所有信息集都是单元素集合,所以每个节点都生成一个子博弈。在每个节点的优化生成所有子博弈的 Nash 均衡。所以,利用反向归纳法得到的解都是 SPNE。

7.3 单偏离原则

有限博弈中的子博弈精炼有个相当有用的特征。在验证某个策略组合 $s(\cdot)$ 是否是子博弈精炼时,只需验证没有参与者有动机在任何一个信息集上偏离它。我们不需考虑在多个信息集上的多个偏离。这一特征被称为单偏离原则,其背后的原理如下:如果在某个有利可图的偏离中,涉及多个历史,那么在其中最后一个历史上的偏离,在自此历史开始的子博弈中有利可图。就是说,如果某个策略组合不是子博弈精炼,可能就存在很多有利可图的偏离。其中至少一个偏离仅为单时期偏离。

作为习题,请你在单人扩展式博弈背景中给出这一原理的直接证明。我们现在就一般有限博弈,给出这一原理的正式陈述和证明。陈述这一原理的一种方法是定义单阶段偏离。设 H 是扩展式博弈中非终点历史构成的集合, h' 是 H 的元素。给定策略 $s(\cdot)$,如果对所有 $h \in H \setminus \{h'\}$,有 $s(h) = s^{h'}(h)$,则策略 $s^{h'}(\cdot)$ 中包含单偏离。

定理 7.2 (单偏离原则) 给定一个有限扩展式博弈,设它的非终点历史集为 H 策略组合,则 $s(\cdot)$ 是子博弈精炼 Nash 均衡当且仅当对每个 $h' \in H$, $p(h')$ 从 $s(\cdot)$ 中获得的收益至少和他从单偏离 $s^{h'}(\cdot)$ 中获得的收益相同。

证明: 由于每个 SPNE 是 Nash 均衡,所以没有人有动机单方面改变自己的策略。就是说,对每个 $h' \in H$, $p(h')$ 从 $s(\cdot)$ 中得到的收益至少和他从 $s^{h'}(\cdot)$ 中得到的收益相同。也就是说,在任何 SPNE 中,没有任何的单阶段偏离是有利可图的。再来证明单偏离原则的充分性。设存在策略组合 $s(\cdot)$,它不是 SPNE,但是其中也不存在有利可图的单偏离。由于 $s(\cdot)$ 不是 SPNE,所以存在某个子博弈, $s(\cdot)$ 在此子博弈上的限制不是 Nash 均衡。设 H' 为构成这一子博弈的最小历史集之一。由于所探讨的是有限博弈,所以存在这样的集合。设 $h^1 \in H'$ 是这个子博弈的初始节点, $i = p(h^1)$ 是在节点 h^1 行动的参与者。通过构造,对于在历史 h^1 选择偏离的参与者 i ,存在对 $s(\cdot)$ 在这一子博弈上的限制的有利偏离。由于我们已经假定任何单偏离都无利可图,所以在 h^1 之后肯定存在非终点历史 h^2 ,在 h^1 和 h^2 的偏离能带给 i 比均衡点更高的收益。但这意味着在 h^2 开始的子博弈中,参与者 i 有有利的偏离。而这显然与前提,即 H' 是 $s(\cdot)$ 的限制不构成 Nash 均衡的最小历史集之一,相矛盾。□

只要连续性假设得到满足,单偏离原则也适用于无限博弈。在重复博弈一章中,我们将继续讨论这一命题。

7.4 子博弈精炼和精炼均衡

子博弈精炼虽然一般用于求解扩展式博弈,但它与第 5 章所讨论的精炼概念有密切联系。前面说过,一个策略组合是精炼的,如果它是在参与者被迫使用完全混合策略时,它是构成最优反应的混合策略组合列的极限。我们用革命博弈的标准形式说明精炼和子博弈精炼之间的关系。我们来看 Nash 均衡组合 $(C, S$ 和 $T)$ 。

假设每一方都必须对其每个纯策略赋予至少为 ϵ 的概率。这就是说,在使用其所偏好的策略时,参与者可能犯错误。

我们现在证明,在 A 国以概率 ϵ 使用其两个策略时, S 和 T 不可能是 B 国的最优反应。如果 A 以概率 ϵ 选择 R , B 从 S 和 T 中得到的期望效用为 $\epsilon(4p - 6 - 2(1 - p)) + (1 - \epsilon)6$, 而从 G 和 T 得到的期望效用为 $(1 - \epsilon)6$ 。由于 $4p - 6 - 2(1 - p) < 0$, B 会以最大概率使用 G 和 T 。所以,在 ϵ 收敛于 0 时, $(C, S$ 和 $T)$ 不可能是完全混合 Nash 均衡的极限。就是说, $(C, S$ 和 $T)$ 不是 SPNE, 也不是标准式博弈中的精炼均衡。作为习题,请你证明只有 $(R, G$ 和 $T)$ 是精炼均衡。

所有的子博弈精炼 Nash 均衡都是精炼均衡,但是,有的精炼均衡并不是子博弈精炼均衡。之所以有这种问题,是因为表示为标准形式的扩展式博弈,在同一个参与者在扩展形式中不止行动一次时,常常产生颤抖上的相关性。解决这一问题的一种方法是,重新定义参与者,使得任何人不会在一个以上的信息集上行动。这种方法产生的博弈形势有时被称为博弈的代理人形式。在这种表述形式中,精炼均衡集和子博弈精炼均衡集相一致。我们把这一命题的证明留做习题。

7.5 应用:议程控制

139

7.5.1 Romer-Rosenthal 模型

在美国许多地方,当地学校预算必须由选民投票通过。一般情况下,只有学校董事会有权提出预算案并交付表决。因此,校董会在学校支出方案上拥有唯一的议程控制权。Romer 和 Rosenthal(1978)最早用模型分析了这种形式的议程控制。

我们介绍 Romer-Rosenthal 模型。设 $s \in [0, \infty)$ 表示支出水平, q 表示当前支出水平。校董会想最大化支出水平,所以它的效用函数 $u_b(s)$ 在 s 上严格递增。在博弈的第一阶段,校董会提议支出水平为 s 。这一提案被提交表决,市民(数量为奇数)在多数通过规则下投票。假设所有市民都参加投票,每个人的策略集为 $\{Y, N\}$ 。如果多数人选择 Y , 则 s 成为了新的支出水平;如果多数人选择 N , 则依然实行当前支出水平 q 。投票者在学校支出上有单峰对称偏好,用效用函数 $u_i(s)$ 表示。设 v_i 是第 i 个人的理想点。和第 2 章一样,这种偏好的形式为 $u_i(s) = h(-|s - v_i|)$, 其中 $h(\cdot)$ 是严格递增函数。

我们想找出子博弈精炼均衡。由于博弈的最后一个阶段是多数通过规则下的投票博弈,所以在这最后一个阶段中,存在着支持接受或支持拒绝任意 s 的 Nash 均衡。因此,我们假定投票者在任何投票子博弈中都不使用弱劣策略。给定支出水平 s , 如果 $u_i(s) \geq u_i(q)$, 则投票者 i 选择 Y 。单峰性意味着中间投票人如果偏好 s , 则支出提案在弱非劣投票中通过。中间投票人如果偏好 q , 则这一提案在弱非劣投票中通不过。

给定投票子博弈的均衡,校董会的最优反应是选择可被中间投票人接受的最

大的 s 。设 v_m 是中间投票人的理想点。给定 v_m 和 q , 我们可以算出中间投票人认为优于 q 的支出水平。这样的支出水平 s 满足条件 $v_m(s) \geq v_m(q)$ 。由于 h 是增函数, 所以, 这一不等式意味着:

$$|s - v_m| \leq |q - v_m|$$

因此, 如果 $q < v_m$, 则对所有的 $s \in [q, 2v_m - q]$, 这一不等式成立。如果 $q > v_m$, 提案要得到通过, 要求 $s \in [2v_m - q, q]$ 。就是说, 所能得到的最高预算为 $\max\{q, 2v_m - q\}$ 。校董会想最大化 s , 所以会选择 $\max\{q, 2v_m - q\}$ 。在图 7.8 中, s^* 的均衡值是 v_m 和 q 的函数。

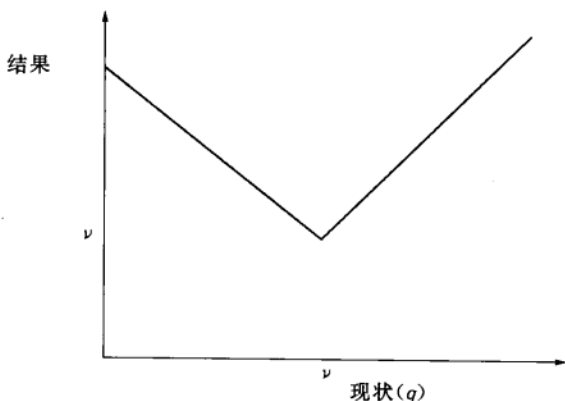


图 7.8 Romer-Rosenthal 博弈中的均衡政策

这个简单的模型就投票人的偏好、当前支出水平和政策结果之间的关系, 得出了一些明确结论。首先, 在当前支出水平很低和中间投票人希望支出数量超过当前水平时, 校董会利用自己对议程的控制权抬高支出水平。由于在当前支出水平很低时, 投票人拒绝高支出方案的威胁不可信, 所以校董会能够实现更高支出。另一个重要的结论是, 支出结果虽然对投票人偏好变化反应敏感(至少在 $q < v_m$ 时), 但是, 支出的增长速度两倍于中间投票人所偏好的支出水平。

7.5.2 总统否决权

在美国和在许多总统制国家中, 总统对立法机构通过的法案拥有否决权。研究总统制的学者常常利用 Romer-Rosenthal 模型, 探讨否决权如何强化总统对立法机构的影响。

为使分析尽可能简单, 我们把立法机构看作是一个行为主体 L , 它在单维空间中有单峰对称偏好。它的理想点是 l 。因此, 对于政策结果 $x \in \mathbb{R}$, 它的效用函数为 $u_l(x) = h(-|x - l|)$, 其中 $h(\cdot)$ 是严格递增函数。总统的理想点为 p , 效用函数为 $u_p(x) = h(-|x - p|)$ 。

这一博弈很简单。在第一个阶段, L 通过法案 b 以改变当前政策 q 。在第二个

阶段,总统 P 决定接受还是否决法案 b 。如果否决,结果便是现状 q 。这里不考虑立法机构推翻总统否決的能力。

我们用反向归纳法求解这一博弈。在最后阶段,只要 $u_p(b) \geq u_p(q)$,即 $-|b-p| \geq -|q-p|$,总统的最优反应是接受这一法案。就是说,如果 $p > q$,他会接受任何法案 b ,只要 $b \in [q, 2p-q]$ 。如果 $p < q$,他会接受任何法案 b ,只要 $b \in [2p-q, q]$ 。设 $P(q)$ 为相对于现状 q ,能被总统接受的法案构成的集合。现在,回到立法机构的决策节点上。立法机构知道哪些法案能被总统接受,所以,它会从 $P(q)$ 中选择自己最偏好的法案。如果 $l \in P(q)$,显然 $b^* = l$ 。如果 $l < \min P(q)$,则 $b^* = \min P(q)$;如果 $l > \max P(q)$,则 $b^* = \max P(q)$ 。

如果 $l > p$, 根据上面的分析,均衡政策结果为:

$$b^* = \begin{cases} 2p-q, & \text{如果 } p > q \text{ 且 } l > 2p-q \\ l, & \text{如果 } p > q \text{ 且 } l < 2p-q \\ l, & \text{如果 } l < q \\ q, & \text{如果 } l > q > p \end{cases}$$

如果 $p > l$, 均衡结果为:

$$b^* = \begin{cases} 2p-q, & \text{如果 } p < q \text{ 且 } l < 2p-q \\ l, & \text{如果 } p < q \text{ 且 } l > 2p-q \\ l, & \text{如果 } l > q \\ q, & \text{如果 } p > q > l \end{cases}$$

在图 7.9 中,均衡结果是 l 、 p 和 q 的函数。这一模型的比较静态特征类似原始 Romer-Rosenthal 模型。具体地说,在现状政策远离总统的理想点时,立法机构能得到更好的结果。另一重要结论是,否決权所产生的影响并不大。在否決行为产生影响(即 $b^* \neq l$)的所有情况中,总统在均衡法案和现状之间是无差异的。最后,这个模型是完美信息博弈,立法机构能够完全预测到总统的行为,所以在均衡

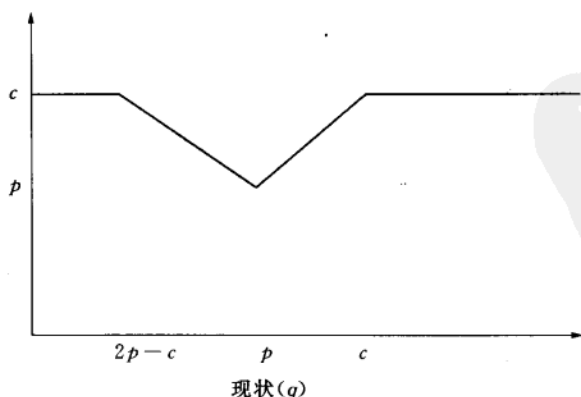


图 7.9 否決博弈中的均衡结果

中,法案不会被否决。在均衡路径上,虽然没有否决行为,但是法案被否决的可能性依然相当重要。在后面几章中,我们将看到,否决行为是均衡策略的组成部分。

7.5.3 推翻总统否决令

在上一个模型中,我们假定总统如果行使否决权,则结果便是当前在推行的政策。这里放弃这一假设,而假定立法机构能够在多数通过的规则下,推翻总统的否决令。立法机构有 n 个成员(n 为奇数)。要推翻总统的否决令,需要有 $k > \frac{n+1}{2}$ 的票数支持。每个立法人员都有单峰偏好,效用函数为 $u_i(x) = h(-|x - l_i|)$, 其中理想点 l_i 是顺序排列的,就是说, $l_i > l_j$ 当且仅当 $i > j$ 。我们进一步假定法案的提出者是处于中间位置的立法人员,其理想点是 $m = l_{\frac{n+1}{2}}$ 。如果法案的提出是通过公开规则(open rule)过程进行的,就会出现这种情况。我们只探讨在所有子博弈中均衡投票策略为弱非劣策略的均衡。

在这些假设下,立法机构要成功推翻总统否决令,需要 $u_k(b) \geq u_k(q)$ 且 $u_{n-k+1}(b) \geq u_{n-k+1}(q)$ 。我们来证明这一点。设 $u_k(b) \geq u_k(q)$ 但 $u_{n-k+1}(b) < u_{n-k+1}(q)$ 。由于所有人的偏好都是单峰的,所以存在某个 $i \in [n-k+1, k]$ 使得对理想点小于 l_i 的所有的立法人员, $u_i(b) < u_i(q)$ 。因此,支持推翻总统否决令的立法人员的数量肯定严格小于 k 。其他可能性背后的逻辑亦如此。由于立法人员 $n-k+1$ 和 k 的支持既是必要的也是充分的,所以他们一般被称为推翻否决令的关键投票人。

由于只在总统行使了否决权时,才可能需要推翻否决令,所以只有一位推翻否决令的关键投票人在策略上重要。设被否决的法案的特征是 $u_p(b) < u_p(q)$ 且 $u_m(b) > u_m(q)$ 。如果 $p < m$, 单峰性和 $l_k > m$ 意味着 $u_k(b) > u_k(q)$ 。于是,是否能推翻否决令只决定于 $n-k+1$ 的偏好。如果 $p > m$, 单峰性和 $l_{n-k+1} < m$ 意味着 $u_{n-k+1}(b) > u_{n-k+1}(q)$ 。就是说,能否推翻总统的否决令,决定于和总统位于中间位置同一侧的关键人的偏好。

根据上述分析,提案人要使提案成功变为政策,必须得到总统的支持,或者得到与总统位于中间位置同一侧的关键投票人的支持。我们来分析提案人如何最优地选择自己的提案。在 $l_{n-k+1} < p < m$ 时,根据单峰性,相对于 q , l_{n-k+1} 和 m 所偏好的任何方案同时也是 p 所偏好的方案。因此,提案人并不需要 l_{n-k+1} 的支持。同样的,在 $p < l_{n-k+1} < m$ 时,提案人只需得到 $n-k+1$ 的支持。 $p > m$ 的情况与此对称。就是说,在总统和与总统位于中间位置同一侧的推翻否决令的关键投票人之间,提案人只需得到最接近自己的那个人的支持。具有这一特征的人,可以称其为关键人。

因此,如果 $p < m$, 关键人的理想点是 $v = \max\{l_{n-k+1}, p\}$; 如果 $p \geq m$, $v = \max\{l_k, p\}$ 。于是,在 $v > m$ 时, SPNE 法案为:

$$b^* = \begin{cases} 2v - q, & \text{如果 } v > q \text{ 且 } m > 2v - q \\ m, & \text{如果 } v > q \text{ 且 } m < 2v - q \\ m, & \text{如果 } m < q \\ q, & \text{如果 } m > q > v \end{cases}$$

在 $v \leq m$ 时, SPNE 法案为:

$$b^* = \begin{cases} 2p - q, & \text{如果 } p < q \text{ 且 } m < 2p - q \\ m, & \text{如果 } p < q \text{ 且 } m > 2p - q \\ m, & \text{如果 } m > q \\ q, & \text{如果 } m > q > p \end{cases}$$

图 7.10 中给出了 k 取不同值时,所对应的均衡结果。随着推翻否决令所需的票数下降,否决权的威力在下降。

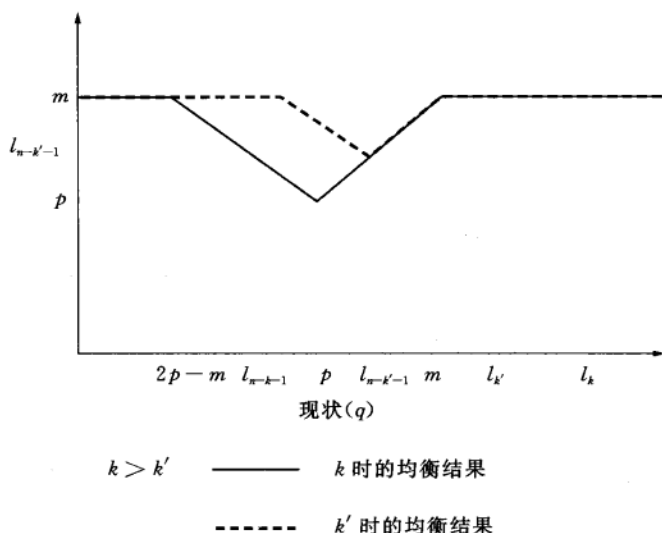


图 7.10 推翻否决令的效应

7.6 应用:力量结构变化模型

Powell(1999)分析了国际体系中力量的巨大变化如何导致武力冲突。设有国家 A 和 B。A 国对 B 国控制的某个地区提出权益要求。这个地区的总价值在每个阶段中始终为 1 元。我们探讨一个两阶段博弈模型,其中,每个国家对每个阶段的结果有相同的价值认定。

首先看 B 国的行为。它可以在第一个阶段中向 A 国提供此地区产出的某一份额,如 $0 \leq x_1 \leq 1$, 以此与 A 国和平共处;它也可以进攻 A 国,武力解决冲突。

如果 B 国选择进攻策略并赢得战争, A 国就放弃自己的权益要求, 博弈结束。B 国如果输掉战争, A 国也就绝对控制了这一地区, 博弈结束。在第一个阶段中如果发生了战争, 战胜国在每个阶段都获得 1 元钱的产出, 战败国则在每个阶段中得到 0 元钱的产出。战胜国和战败国都得付出战争成本 c 。如果 B 国提出了收益分享方案 x_1 , A 国或者接受 x_1 或者拒绝并开战。如果 B 国发起战争, A 国别无选择, 只能迎战。如果在第一个阶段中没有发生战争, 则在第二个阶段中, B 国面临着相同的选择: 向 A 国提供某一产出份额 $0 \leq x_2 \leq 1$, 或者发起战争。如果在第二个阶段中, B 国提出了分享方案, A 国或者接受或者拒绝并对 B 国开战。在第二个阶段中如果发生战争, 战胜国得到 1 元钱的产出, 而战败国得到 0 元钱的产出。第二个阶段中的战争成本对双方来说依然为 c 。

这个模型的一个重要特征是, 随着时间推移, A 国的军力相对于 B 国在加强。在第一个阶段, A 国以概率 p_1 赢得战争; 在第二个阶段, A 国以概率 $p_2 > p_1$ 赢得战争。为使分析有意义, 假设 $p_2 > c > p_1$ 。

就武力冲突发生与否而言, 存在两类均衡: 在第一类均衡中, B 在两个阶段里都向 A 示好; 在第二类均衡中, B 在第一个阶段向 A 发动战争。

假设 $2c > p_2 - 2p_1$ 。在 SPNE 中, B 国在第二个阶段向 A 国提供产出 $p_2 - c$, 在第一个阶段提供产出 $\max\{0, 2p_1 - p_2\}$ 。由于这是个完美信息博弈, 所以我们可以利用反向归纳法证明这些策略是 SPNE 的构成部分。在第二个阶段里, A 从战争中所获得的期望效用是 $p_2 - c$, 所以, B 国至少需要向 A 国提供这一数量的产出才能够防止武力冲突。B 国的纳贡行为给自己实现的收益为 $1 - p_2 + c$ 。由于 B 国从战争中所获得的期望效用是 $1 - p_2 - c$, 所以它所偏好的选择是向 A 国示好。因此, B 国在第二个阶段中的最优行动是向 A 国提供产出 $x_2^* = p_2 - c$ 。

再看第一阶段。A 国如果赢得战争, 则在每个阶段得到 1 元钱收益; 如果输掉战争, 得到 0 元钱的收益。因此, 它从战争中所得到的期望效用是 $2p_1 - c$ 。为向 A 国示好, B 国对 x_1 的选择必须满足 $x_1 + p_2 - c \geq 2p_1 - c$, 即满足 $x_1 \geq 2p_1 - p_2$ 。由于 B 的理性选择是提供满足这一条件的最低数量的 x_1 , 所以, 它的均衡选择是 $x_1^* = \max\{0, 2p_1 - p_2\}$ 。我们现在只需验证 B 国是否愿意支付 $x_1^* + x_2^*$ 。B 国从战争中所得到的期望效用是 $2(1 - p_1) - c$, 所以它愿意向 A 国支付 x_1^* 当且仅当 $1 - x_1^* + 1 - x_2^* \geq 2(1 - p_1) - c$ 。这一条件等价于 $2 - x_1^* - p_2 + c \geq 2(1 - p_1) - c$, 即 $2c \geq p_2 - 2p_1 + x_1^*$ 。如果 $2c > p_2 - 2p_1$, 这一条件得到满足。

假设这一条件没有得到满足, 就是说, $p_2 - 2p_1 > 2c$ 。如果 $c > 0$, 这就意味着 $x_1^* = 0$ 。于是, B 国更愿意发动战争而不是向 A 国纳贡, 因为 $p_2 - 2p_1 + x_1^* > 2c$ 。因此, 当前述条件得不到满足时, 在 SPNE 中 B 国在第一个阶段进攻 A 国。

只有在 p_2 远大于 p_1 , 就是说, 当 A 国在第一个阶段相对于第二个阶段弱很多时, 两个国家和平相处的必要条件才得不到满足。在 A 国弱小时, B 国偏好进攻, 借此避免在 A 国变得强大时, 向 A 国纳贡很多。如果力量的分布是稳定的 (如在 $p_1 = p_2$ 时), 则在均衡中没有战争。

7.7 应用：向民主制度的转轨

Acemoglu 和 Robinson(2005)设计了许多模型探讨在什么条件下，专制政府会推行民主制度。我们分析其中一个模型，借此勾勒其理论框架。

设有两类参与者：富人和穷人。 $\lambda > \frac{1}{2}$ 是居民中穷人的比例， $1-\lambda$ 是居民中富人的比例。每个富人有收入 y^r ，每个穷人有收入 y^p 。社会中的平均收入是 $\bar{y} = \lambda y^p + (1-\lambda)y^r$ 。显然， $y^r > \bar{y} > y^p$ 。在 Acemoglu 和 Robinson 的分析中，一个重要的参数是 θ ，表示全部收入中穷人所占有的份额。它由下面两个恒等式定义：

$$y^p = \frac{\theta \bar{y}}{\lambda} \text{ 和 } y^r = \frac{(1-\theta)\bar{y}}{1-\lambda}$$

就是说， θ 的上升意味着收入差距缩小。

在这个经济中，主要的政策工具是线性的税收或转移支付方案：政府制定比例税率 τ ，然后将税收收入以按人头补贴的形式返回给居民。由于人们有不同的收入，所以他们对税率有不同的偏好。给定税率 τ ，人均税收为 $\tau \bar{y}$ 。为反映所得税的扭曲效应，Acemoglu 和 Robinson 假定实际税收收入小于应征税收总额，因为征税过程中要发生交易成本。在应征税收为 $\tau \bar{y}$ 时，交易成本为 $C(\tau)\bar{y}$ ， $C(\tau)$ 是凸的和递增的。为使分析简单并得到闭式解，设 $C(\tau) = \frac{1}{2}\tau^2$ 。在扣减了这部分净损耗后，可用于转移支付的资金总额为 $T = (\tau - C(\tau))\bar{y}$ ；在纳税或得到转移支付后，类型为 $i \in \{r, p\}$ 的居民的收入为：

$$V^i(\tau) = (1-\tau)y^i + \left(\tau - \frac{1}{2}\tau^2\right)\bar{y}$$

我们来分析富人和穷人各自所偏好的税率。把税后收入对税率求导并让导数等于 0，就得到对穷人最优的税率所满足的一阶条件：

$$\bar{y} - y^p - \tau \bar{y} = 0$$

代入恒等式 $y^p = \frac{\theta \bar{y}}{\lambda}$ 中，穷人所偏好的税率为：

$$\tau^p = \frac{\lambda - \theta}{\lambda}$$

由于穷人的收入比富人低，他们在社会收入中所占有的份额低于其在人口中的比例，所以， $\lambda > \theta$ 。因此， $0 < \tau^p < 1$ 。同时， τ^p 在 θ 上递减，就是说，穷人所偏好的税率随收入差距而增大。

再分析富人的偏好。最优税率如果为内解，则一阶必要条件为：

$$\bar{y} - y^r - \tau \bar{y} = 0$$

只有在税率为负时,这一条件才得到满足。由于负税率不可行,所以,富人最偏好的可行税率为 $\tau^* = 0$ 。

在 Acemoglu 和 Robinson 模型中,在每个时期里都存在政治冲击,这种政治冲击可能会使现有体制被推翻,取而代之的是左翼专制政权。在冲击为 S 时, $1 - \mu^S$ 比例的社会收入被破坏。其中, $S \in \{H, L\}$, $\mu^H > \mu^L$ 。就是说,在状态 H ,推翻现有体制的成本要比状态 L 中的成本低。在革命中,富人的收入被没收,平均分配给穷人。所以,在状态 S 中,穷人在革命后的收入为:

$$V^p(R, \mu^S) = \frac{\mu^S \bar{y}}{\lambda}$$

为简化分析,他们假定在革命后,富人没有任何收入,即 $V^r(R, \mu^S) = 0$ 。

为说明 Acemoglu 和 Robinson 模型的结论,我们探讨革命在多大程度上构成实实在在的威胁。如果在税率为 0 时,就是说,在政府推行的是富人的理想税率时,穷人在革命和专制之间偏好革命时,革命约束条件在状态 S 中就是紧的。在 $V^p(R, \mu^S) > V^p(0)$ 时,便是这种情况。代入前述表达式,我们就能看到,在 $\mu^S > \theta$ 时这一约束条件是紧的。

给定这样的经济结构和革命成本,我们探讨 Acemoglu 和 Robinson 所研究的一个扩展式博弈。图 7.11 给出了这一博弈的一个简化博弈树——只给出了在状态 S 之后的部分。

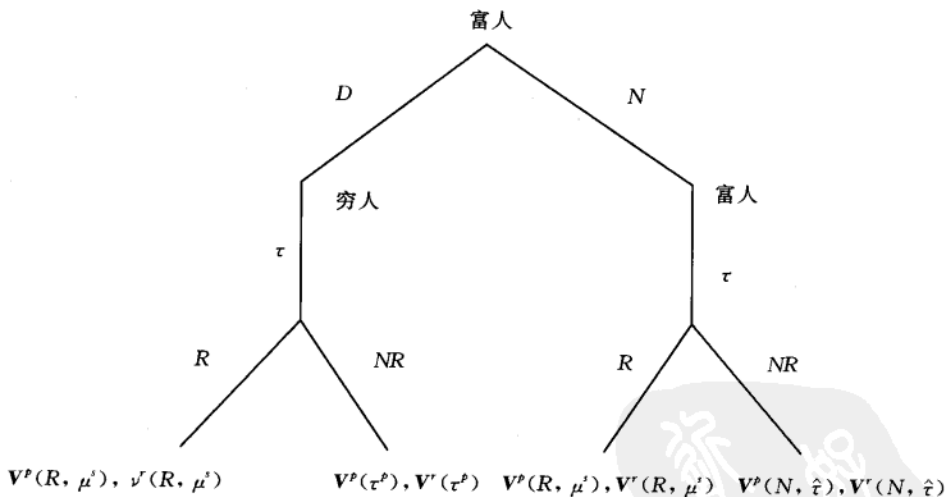


图 7.11 民主化博弈

首先,状态 H 或 L 被揭示。然后,富人首先行动,选择“过渡到民主政体”(D)还是依然“实行专制控制”(N)。富人如果选择了 N ,则同时选择税率 $\hat{\tau}$ 。

在富人做出选择后,穷人决定是“发起革命”(R)还是“接受富人的决策”(NR)。如果革命,穷人和富人的收益分别为 $V^p(R, \mu^S) = \frac{\mu^S \bar{y}}{\lambda}$ 和 $V^r(R, \mu^S) = 0$ 。在富人

选择 D 时,如果穷人不革命,税率就由多数通过的投票规则决定。由于中间选民是穷人,所以均衡税率是 τ^p ,由 D 所实现的收益分别为 $V^p(\tau^p)$ 和 $V^r(\tau^p)$ 。比较 R 和 D 带来的收益,很容易看出,在状态 S 中,穷人偏好革命而不接受民主政体当且仅当 $\mu^s > \theta + \tau^p(\lambda - \theta) - \frac{1}{2}\tau^p\lambda$ 。这个充分必要条件可以简化为 $\mu^s > \theta + \frac{(\lambda - \theta)^2}{2\lambda}$ 。

假设富人选择 N 而穷人不想革命。Acemoglu 和 Robinson 认为,在革命威胁消除后,富人可能无法保持税率为 $\hat{\tau} > 0$ 。为分析这一问题,他们假定,富人保持税率 $\hat{\tau}$ 的概率为 p ,而撕毁承诺转而选择 $\tau^r = 0$ 的概率为 $1 - p$ 。给定富人对税率的初始选择和撕毁税率承诺的可能性, N 带给双方的效用分别为:

$$\begin{aligned} V^p(N, \hat{\tau}) &= (1-p)y^p + p\left[(1-\tau)y^p + \left(\tau - \frac{1}{2}\tau^2\right)\bar{y}\right] \\ &= y^p + p\left[\tau(\bar{y} - y^p) - \frac{1}{2}\tau^2\bar{y}\right] \\ V^r(N, \hat{\tau}) &= (1-p)y^r + p\left[(1-\tau)y^r + \left(\tau - \frac{1}{2}\tau^2\right)\bar{y}\right] \\ &= y^r + p\left[\tau(\bar{y} - y^r) - \frac{1}{2}\tau^2\bar{y}\right] \end{aligned}$$

给定这些收益,容易看出,如果 $\mu^s > \theta + p\left[\hat{\tau}(\lambda - \theta) - \frac{1}{2}\hat{\tau}^2\lambda\right]$,穷人会通过革命反对 N 。

为减少所需研究的情况的数量,Acemoglu 和 Robinson 假定 $\mu^L < \theta$ 。于是,在状态 L 中,穷人永远不会选择革命。我们就有三种可能性需要探讨:

(1) $\mu^H < \theta$ 。在任何情况中,革命的约束条件都不是紧的。于是,唯一的 SPNE 由 N 、零税率和不革命构成。

(2) $\mu^H > \theta + \frac{(\lambda - \theta)^2}{2\lambda}$ 。此时,甚至民主政治都无法阻止穷人革命,革命必然发生。

(3) $\theta + \frac{(\lambda - \theta)^2}{2\lambda} > \mu^H > \theta$ 。此时,富人有可能为防止革命推行税率 $\hat{\tau}$ 。从前面的分析中我们知道,要做到这一点,富人所设定的税率必须满足:

$$p < \frac{\mu^H - \theta}{\hat{\tau}(\lambda - \theta) - \frac{1}{2}\hat{\tau}^2\lambda}$$

但是,如果:

$$p < \frac{\mu^H - \theta}{\tau^p(\lambda - \theta) - \frac{1}{2}\hat{\tau}^p\lambda} = \frac{2\lambda(\mu^H - \theta)}{(\lambda - \theta)^2}$$

富人就会选择 D 而不会把税率定在高于 τ^p 的水平上。于是,存在临界值:

$$p^* = \frac{2\lambda(\mu^H - \theta)}{(\lambda - \theta)^2}$$

使得在 $p^* > p$ 时, 民主政治便是最终结果。就是说, 当富人无法承诺保持高税率时, 他们通过转变为民主政治而防止革命。

为知道在什么条件下可能发生向民主的转轨, 我们可以分析参数上的变化如何影响 p^* 。不难看出, p^* 在 μ^H 上递增。这意味着, 在革命成本很低时, 富人更可能支持民主化。其次, p^* 和民主化的发生概率在 θ 上递减, 因为收入差距扩大使得革命成为了对穷人更有吸引力的选择, 而这又意味着富人为防止革命必须在税率上做出更大的让步。如果富人很难做出这种让步, 向民主的转轨就会发生。

7.8 应用: 政党联盟的形成模型

政治博弈论最早的应用之一是研究联盟的形成(Riker, 1962)。早期的模型建立在合作博弈论和社会选择框架内, 但后来的大量研究使用的是非合作讨价还价模型。

在本节中, 我们介绍 Austen-Smith 和 Banks(1998)的联合执政模型。有三个政党 α 、 β 和 γ , 它们的政策立场 p_α 、 p_β 和 p_γ 是公开信息, 来自单维度政策空间 $P \subset \mathbb{R}$ 且 $p_\alpha > p_\beta > p_\gamma$ 。向量 $\omega = (\omega_\alpha, \omega_\beta, \omega_\gamma)$ 是在最近一轮选举中, 三个政党的得票份额。设所有党派的得票份额都低于 $\frac{1}{2}$, 所以, 政府只能实行多政党联合执政。为简化分析, 假定这些票数份额是外生参数, 不过在 Austen-Smith 和 Banks 模型中, 它们内生地产生自一个投票模型。如果 $C \subset \Omega = \{\alpha, \beta, \gamma\}$ 是个政党联盟, 则这一联盟得到的票数份额为:

$$\omega_C = \sum_{k \in C} \omega_k$$

如果 $\omega_C > \frac{1}{2}$, 则联盟 C 胜出; 由于我们假定任何政党的得票份额都小于 $\frac{1}{2}$, 所以任何联盟要胜出, 至少必须由两个政党组成。

三个党派就新一届政府的组成而讨价还价。在这一过程中, 它们选择政策 $y \in P$, 并在各政党间分配政府职位。假设政府职位集 G 固定不变。与 Austen-Smith 和 Banks 一样, 我们假定 G 是无限可分的, 任何职位分配方案 $g = \{g_\alpha, g_\beta, g_\gamma\}$ 满足 $\sum_{k \in C} g_k = G$ 。

每个党派都有可分离的效用函数, 其中, 政策带来二次方收益, 职位数量带来线性可加的收益。于是, 党派 k 从政策 y 和职位分配 g 中获得的收益为:

$$-(y - p_k)^2 + g_k$$

党派之间的讨价还价过程如下: 首先, 得票最多的政党, 譬如说政党 k , 被指定组建执政联盟。政党 k 提出政策 y_k 和分配方案 g_k 。如果它的联盟伙伴接受, 就执行 y_k 和 g_k , 博弈结束。如果它的某个联盟伙伴拒绝, 则由得票多的党派臂

如说党派 l 组建执政联盟。它提出 y_l 和 g_l 。如果它组阁失败,则由得票最少的政党组建政府。如果它也失败,则由看守政府负责政府事务,推行当前政策 p_q 并选择 $g_c = \{0, 0, 0\}$ 。

我们用反向归纳法求解这一博弈。在每个组阁阶段,设组阁党派为 $j \in \{\alpha, \beta, \gamma, c\}$, v_i^j 是党派 i 的效用, $i \in \{\alpha, \beta, \gamma\}$ 。看守政府带给党派 i 的效用为 $v_i^c = -(p_q - p_k)^2$ 。为简化分析,我们假定对所有的 j 和 k , $v_i^j < -(p_j - p_k)^2$, 就是说,相对于看守政府,每个党派 k 都偏好党派 j 的理想点,即使自己在政府中没有任何职位。正式地说,这要求 $p_q \notin [2p_\gamma - p_\alpha, 2p_\alpha - p_\gamma]$ 。

假设投票份额结构为 $\omega_\alpha > \omega_\beta > \omega_\gamma$ 。则在第三个阶段,政党 γ 负责组阁。根据前述假设,所有党派都偏好 $y_\gamma = p_\gamma$ 和 $g_\gamma = (0, 0, G)$ 而不愿由看守政府执政,所以, (y_γ, g_γ) 是党派 γ 的最优选择。

在第二个阶段,党派 β 负责组阁。拒绝 β 的组阁方案从而进入到 γ 的组阁阶段给另两个政党带来的效用分别为 $v_\gamma^\beta = G$ 和 $v_\alpha^\beta = -(p_\gamma - p_\alpha)^2$ 。由于 γ 从反对 β 的组阁方案中得到最大效用,所以 β 只能求助党派 α 。由于相对于 γ 的组阁方案, α 偏好 $y_\beta = p_\beta$ 和 $g_\beta = (0, G, 0)$, 所以 α 会接受 β 的理想点并且不要求任何政府职位。

在博弈的第一阶段中, α 提出组阁方案。拒绝 α 的组阁方案而进入到 β 的组阁阶段给另两个政党带来的效用分别为 $v_\beta^\alpha = G$ 和 $v_\gamma^\alpha = -(p_\beta - p_\gamma)^2$ 。显然, α 没办法与 β 组阁,而只能与 γ 组建联盟。所以, α 会通过对于 y_α 和 g_γ 的选择,在约束条件 $-(y_\alpha - p_\gamma)^2 + g_\gamma \geq -(p_\beta - p_\gamma)^2$ 和 $G \geq g_\gamma \geq 0$ 下,最大化 $-(y_\alpha - p_\alpha)^2 + G - g_\gamma$ 。^① 根据是否存在角解 $g_\gamma^* = G$ 或 $g_\gamma^* = 0$,有以下三种不同情况:

(1) 如果 $p_\alpha - p_\beta \geq p_\beta - p_\gamma$ 且 $G \geq \frac{1}{4}(p_\alpha - p_\gamma)^2 - (p_\beta - p_\gamma)^2$, 则 $y_\alpha^* = \frac{p_\alpha + p_\gamma}{2}$ 且 $g_\gamma^* = \frac{1}{4}(p_\alpha - p_\gamma)^2 - (p_\beta - p_\gamma)^2$ 。

(2) 如果 $p_\alpha - p_\beta \geq p_\beta - p_\gamma$ 且 $G < \frac{1}{4}(p_\alpha - p_\gamma)^2 - (p_\beta - p_\gamma)^2$, 则 $y_\alpha^* = p_\gamma + \sqrt{G + (p_\beta - p_\gamma)^2}$ 且 $g_\gamma^* = G$ 。

(3) 如果 $p_\alpha - p_\beta < p_\beta - p_\gamma$, 则 $y_\alpha^* = p_\beta$ 且 $g_\gamma^* = 0$ 。

在前两种情况中,从 α 的理想点到 p_β 的距离大于从 p_β 到 γ 的理想点的距离。因此,政党 α 愿意放弃政府中的一些职位并在政策上妥协。在 G 足够地大时, α 所提供的妥协政策是 $y_\alpha^* = \frac{p_\alpha + p_\gamma}{2}$, 反映了在其政策目标和它控制政府职位的意愿之间的最优替代。但是,在 G 很小时, α 愿意放弃所有政府职位以使政策靠拢其理想点。在第三种情况中,相对于 γ , α 在 p_β 下得到足够高的效用,所以它不愿意为使政策靠拢自己的理想点而用政府职位补偿 γ 。这个结果的一个有趣的特点

① 书末的数学附录中介绍了约束条件下的优化。

是,这一联盟是两个极端党派之间的非连通的结盟。这一结论与另一种观点——政策导向的政党与其意识形态上的盟友组建联盟(Axelrod, 1970)——截然对立。

假设 $\omega_\beta > \omega_\alpha > \omega_\gamma$ 。在最后一个阶段中, γ 依然选择 $y_\gamma^* = p_\gamma$ 和 $g_\gamma = (0, 0, G)$ 。我们看在第二个阶段中 α 的选择。它显然没办法拉拢 γ , 因此只能设法与 β 结盟。所以, α 会通过 y_α 和 g_β 的选择, 在约束条件 $-(y_\alpha - p_\beta)^2 + g_\beta \geq -(p_\beta - p_\gamma)^2$ 和 $G \geq g_\beta \geq 0$ 下, 最大化 $-(y_\alpha - p_\alpha)^2 + G - g_\beta$ 。这个问题的解导致了以下四种不同情况。

(1) 如果 $\frac{p_\alpha + p_\beta}{2} \geq 2p_\beta - p_\gamma$ 且 $G \geq \frac{1}{4}(p_\alpha - p_\beta)^2 - (p_\beta - p_\gamma)^2$, 则 $y_\alpha^* = \frac{p_\alpha + p_\beta}{2}$ 且 $g_\beta^* = \frac{1}{4}(p_\alpha - p_\beta)^2 - (p_\beta - p_\gamma)^2$ 。

(2) 如果 $\frac{p_\alpha + p_\beta}{2} \geq 2p_\beta - p_\gamma$ 且 $G < \frac{1}{4}(p_\alpha - p_\beta)^2 - (p_\beta - p_\gamma)^2$, 则 $y_\alpha^* = p_\beta + \sqrt{G + (p_\beta - p_\gamma)^2}$ 且 $g_\beta^* = G$ 。

(3) 如果 $p_\alpha > 2p_\beta - p_\gamma > \frac{p_\alpha + p_\beta}{2}$, 则 $y_\alpha^* = 2p_\beta - p_\gamma$ 且 $g_\beta^* = 0$ 。

(4) 如果 $p_\alpha < 2p_\beta - p_\gamma$, 则 $y_\alpha^* = p_\alpha$ 且 $g_\beta^* = 0$ 。

在后两种情况中, α 与 β 之间的讨价还价类似 Romer-Rosenthal 政策博弈, 只是这里的底线是 p_γ 。在第 1 和第 2 种情况中, 政党 α 让出政府职位以求政策向自己的理想点移动, 只是在这里, 政策的移动程度超过了 Romer-Rosenthal 模型中的现状点 $2p_\beta - p_\gamma$ 。

这些情况尽管有些复杂, 但有一点必须注意, 那就是根据这一模型, $y_\alpha^* > p_\beta$ 且 $g_\gamma^* = 0$ 。就是说, 在第一个时期里, 当党派 β 组阁时, 它知道党派 γ 接受 $y_\beta^* = p_\beta$ 和 $g_\gamma^* = 0$ 。因此, 子博弈精炼 Nash 均衡为 $y^* = p_\beta$ 和 $g^* = (0, G, 0)$ 。这种情况中产生的是连通的联盟, 执行的是中间党派理想点。

我们把剩余情况的证明留做习题。表 7.5 概括了所有情况中的结果。

表 7.5 Austen-Smith 和 Banks 模型中的结果

情 况	执政联盟	政 策
$\omega_\alpha > \omega_\beta > \omega_\gamma$	α 和 γ	$\frac{p_\alpha + p_\gamma}{2}$
$\omega_\beta > \omega_\alpha > \omega_\gamma$	β 和 γ	p_β
$\omega_\beta > \omega_\gamma > \omega_\alpha$	β 和 α	p_β
$\omega_\alpha > \omega_\gamma > \omega_\beta$	α 和 β	$\frac{p_\alpha + p_\beta}{2}$
$\omega_\gamma > \omega_\alpha > \omega_\beta$	γ 和 β	$\frac{p_\gamma + p_\beta}{2}$
$\omega_\gamma > \omega_\beta > \omega_\alpha$	γ 和 α	$\frac{p_\alpha + p_\gamma}{2}$

Austen-Smith 和 Banks 模型的一个重要结论是,政府的构成和它所执行的政策决定于组阁顺序,而这一顺序又决定于各政党的得票份额。由于得票份额由投票行为决定,所以,找到模型结论的关键,是了解选民们在预知到政党间的讨价还价的前提下如何投票。我们建议读者阅读 Austen-Smith 和 Banks 的论文,了解他们对投票博弈的分析。

7.9 习题

习题 7.1 将蜈蚣博弈表述为标准式博弈,分析 Nash 均衡集的特征。

习题 7.2 证明在表 7.2 的革命博弈中,策略组合 $(R, G$ 和 $T)$ 是唯一的精炼均衡。

习题 7.3 Diane 在为民主政治研究中心筹措咖啡基金。这个基金如果正常运转,她每个月必须向中心至少三位成员每人收取 2 美元。任何成员的出资额都不会超过 2 美元,并且,Diane 无法拒绝没有出资者饮用咖啡。每个成员从享用咖啡中得到相当于 10 美元的收益。如果募集不到 6 美元,Diane 可以将这笔钱装入自己口袋而不提供咖啡。如果资金数额超过 6 美元,Diane 就得提供咖啡,但可以把多余部分装入自己口袋。

(1) Diane 决定依下列顺序请中心各个成员出资:Arnold、Bartels、Lewis、Prior 和 Romer。每个成员能观察到在自己之前的出资情况。这一博弈的 Nash 均衡集是什么?在反向归纳法下“存活”下来的唯一的 Nash 均衡是什么?

(2) 现在假设 Lewis、Prior 和 Romer 在决定是否出资时,不知道 Arnold 和 Bartels 是否提供了资金;假设 Lewis、Prior 和 Romer 必须同时决定是否出资。为这一博弈的扩展形式画出博弈树并说明博弈中的信息集的特征。这一博弈的子博弈精炼均衡是什么?和前一博弈相比,Diane 现在的收入是改善了还是恶化了?

(3) 现在由 Diane 决定这一博弈的信息结构(就是说,她可以决定在哪个阶段上揭示哪些人提供了资金)。假设她想最大化自己能得到的资金数量而不关心咖啡消费。她应该选择上述哪个博弈?

习题 7.4 分析下面的序贯讨价还价问题,证明如下策略是唯一的子博弈精炼 Nash 均衡:

参与者 1:建议 $x_2 = \delta$, 如果这一方案被拒绝,就接受在下一阶段中对方提出的任何方案。

参与者 2:如果 $x_2 \geq \delta$, 就接受第一阶段中对方提出的方案;否则,拒绝对方的方案并在第二个阶段中提出方案 $x_2 = \delta$ 。

提示:首先假设参与者 1 的策略是,只有当第二个阶段的方案带给她的收益至少为 $\varepsilon > 0$ 时,才接受对方方案。证明这一策略不是自第一阶段的方案被拒绝后开始的子博弈中的 Nash 均衡。

习题 7.5 我们分析两个参与者之间的讨价还价问题。假设参与者 1 向参与

者 2 提出的分配方案为 (x_1^j, x_2^j) , 其中的 x_j^j 是在参与者 i 提出的这一方案被对方 j 接受时, j 所获得的收益。在第一个阶段中, 参与者 1 提出了对 1 元钱的分配方案 (双方的所得皆非负且和为 1)。参与者 2 面对着这一方案, 决定自己是接受还是拒绝。如果接受, 则按照方案分配。如果拒绝, 则由她提出分配 $\delta < 1$ 美元的方案 (双方的所得皆非负且和为 δ)。参与者 1 决定是接受还是拒绝这一方案。如果他接受, 则按照方案进行分配。如果他拒绝, 则双方各得 $\delta/2$ 。设两个参与者都最大化自己的所得, 找出这一博弈的子博弈精炼均衡。

习题 7.6 Groseclose(1999) 设有 N 个立法者, 各自的政策偏好可以表示为 $u_i(x)$, 其中 $x \in \mathbb{R}$ 为某一政策。他们以投票方式在法案 x_B 和现状 x_0 之间做出选择, 设 $x_B > x_0$ 。 $\alpha_i = u_i(x_B) - u_i(x_0)$ 相对于 x_0 , 立法者 i 偏好 x_B 的程度。无论法案是否通过, 每个立法者从这一法案进行的投票行为中获得的政策收益为 α_i 。设有两个人, L 和 R , 在购买选票。他们分别有净偏好参数 α_L 和 α_R 。买票者 L 想阻止 x_B , 因此 $\alpha_L < 0$; 买票者 R 希望这一法案通过, 因此 $\alpha_R > 0$ 。分析下面的模型。 R 首先行动, 向同意投票支持这一法案的每个立法者提供贿赂 z_i^R 。然后, L 行动, 向每个立法者提供贿赂 z_i^L ; 作为交换, 立法者投票反对 x_B 。于是, 投票支持 x_B 的收益是 $\alpha_i + z_i^R$, 而投票反对 x_B 的收益是 $-\alpha_i + z_i^L$ 。

(1) 设 $N = 5$, $\alpha_i = -3 + i$, 其中 $i = 1, \dots, 5$ 。对任意的 α_L 和 α_R , 求子博弈精炼 Nash 均衡。

(2) 最终胜出的一方人数远超过半数吗?

习题 7.7 在一个单人博弈中, 参与者 1 选择 A 或 B 。如果选择 A , 他得接着在 C 和 D 之间做出选择; 如果选择 B , 他得接着在 E 和 F 之间做出选择。设每个路径带来的收益分别为: $u(A, C) = u_1$ 、 $u(A, D) = u_2$ 、 $u(B, E) = u_3$ 和 $u(B, F) = u_4$ 。对任意的 u_1 、 u_2 、 u_3 和 u_4 , 证明单偏离原则适用于这一博弈。具体地说, 证明如果某一策略不是子博弈精炼均衡, 则存在有利可图的单偏离。

习题 7.8(*) 在一个扩展式博弈中, 每个参与者仅在一个信息集上行动 [就是说, 映射 $p(\cdot)$ 是个双射]。再假定每个信息集是单元素集合。证明在这种博弈中, 子博弈精炼 Nash 均衡集与在标准形式中为精炼均衡的 Nash 均衡集相同。

习题 7.9 对 Austen-Smith 和 Banks(1988) 模型中剩余的几种情况, 推导出政策结果和政党联盟。

不完全信息动态博弈

第6章指出,其他参与者偏好上的不确定性从根本上改变了静态标准式博弈的策略环境。在动态多阶段博弈中,不确定性导致了更有意思的策略可能性。在图7.1的革命博弈中,唯一的子博弈精炼均衡是殖民地发动革命并独立($R, (G, T)$)。与图7.1不同,在图8.1的博弈中,B国不用发生成本就能武力镇压A国革命,唯一的子博弈精炼均衡是($C, (S, T)$)。

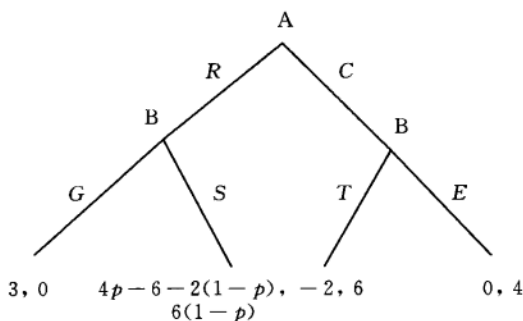


图 8.1 革命博弈

本章要探讨的问题是:在A国不知道两幅图中的哪一幅准确描述了自己所面对的博弈时,它应该怎么做呢?

在我们现在所讨论的博弈中,参与者面对着其他参与者偏好上的不确定性。这种形式的博弈被称为不完全信息博弈。与第6章一样,我们用 Harsanyi 方法处理这种不确定性。其他参与者收益上的不确定性,被转化为参与者不确定自己位于博弈的哪个节点上。这种方法中用到了一个虚拟参与者——自然。自然从某个已知概率分布中随机选择参与者类型。但是,不是所有参与者都能观察到自然的抽取结果。为反映参与者 i 不知道参与者 j 的偏好的事实,我们假定自然在参与者 i 决策之前就选择了参与者 j 的收益(即类型);在参与者 i 所面对的信息集中有多个节点,因为她没有观察到自然做出的选择。这种方法把不完全信息博弈——参与者不知道进行的是哪个博弈——转变为了非完美信息博弈——参与者知道所进行的博弈,但是不知道自己在博

弈中的位置。^①

为更详实地讨论这些问题,我们看两个国家间的冲突模型。A 国首先选择是否挑起冲突。如果没有冲突发生,博弈结束。如果它挑起冲突,B 国决定是顺服还是抵抗。表 8.1 和表 8.2 给出了两种情况中双方从博弈中获得的收益。

表 8.1 博弈 I

A 没有挑起冲突(0, 0)
A 挑起冲突,B 不抵抗(4, -4)
A 挑起冲突,B 抵抗(-8, -8)

表 8.2 博弈 II

A 没有挑起冲突(0, 0)
A 挑起冲突,B 不抵抗(4, -4)
A 挑起冲突,B 抵抗(-8, -3)

如果双方进行博弈 I 的概率是 p , 进行博弈 II 的概率是 $1-p$, 就可能有几种不同的信息结构。

(1) 图 8.2: 两个国家都没有观察到自然所采取的行动。此时我们说, 信息是不完善的但却是对称的, 因为两个国家处于相同情况中。这个博弈容易分析, 我们只需计算 B 国从抵抗中获得的期望效用, 并据此修改博弈即可。B 国从抵抗中获得的期望效用是 $-p \cdot 8 - 3(1-p) = -3 - 5p$, 所以只要 $p < \frac{1}{5}$, 它就会抵抗。就是说, 如果 $p < \frac{1}{5}$, 博弈结果为{不挑起冲突, 抵抗}; 否则, 结果为{挑起冲突, 不抵抗}。

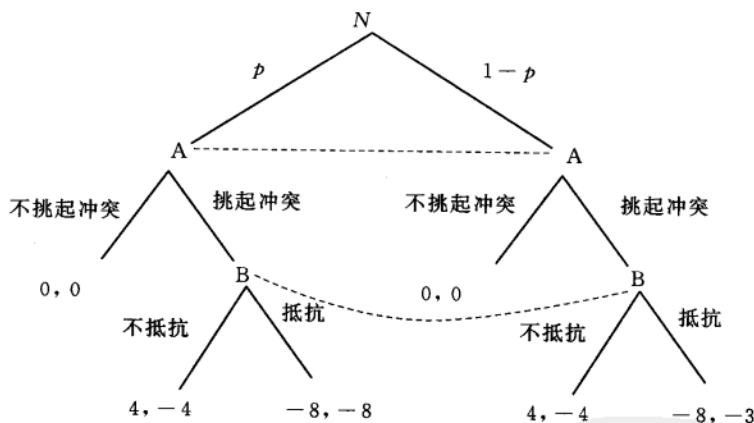


图 8.2 信息结构 I

(2) 图 8.3: 只有 B 国观察到自然的行动。在这一信息结构中, A 国不知道 B 国的选择。由于 B 国只在博弈 II 中才抵抗, 所以, A 国从挑起冲突中获得的期望效用是 $4p - (1-p)8 = -8 + 12p$ 。因此, 只有在 $p > \frac{2}{3}$ 时, A 国才挑起冲突。

① 不完全信息博弈使得均衡概念成为了问题。但是, Mertens 和 Zamir(1985)指出, 在某些技术性条件下, 关于不完全信息的任何描述, 都能概括为贝叶斯博弈。

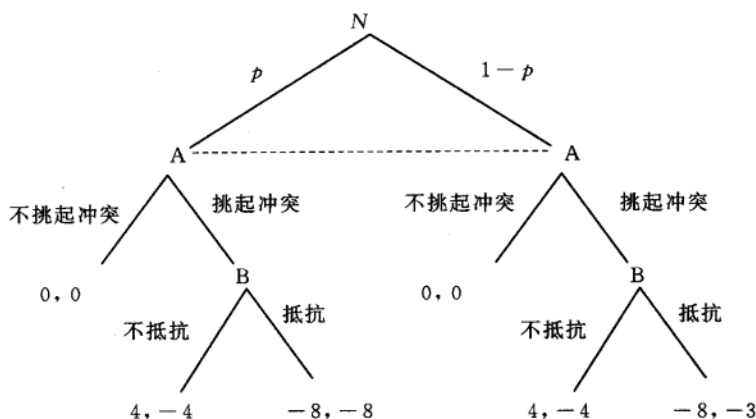


图 8.3 信息结构 II

(3) 图 8.4: 只有 A 国观察到自然的行动。这一博弈有不对称的信息, 其中的策略环境与前两种情况差异很大。B 国不知道自然的选择, 但是他知道 A 国知道自然的选择。所以, A 国肯定考虑自己的行为如何传递出关于自然所做出的选择的信息。为说明这种信息结构影响行为, 我们探讨进行这一博弈的一种方式: A 国在博弈 I 中挑起冲突, 但是在博弈 II 中却不这样做。如果 A 国采取如此策略, B 国就会从 A 国挑起冲突的行为中推断他们在进行博弈 I, 于是采取不抵抗政策。但是, 如果 B 国以这种方式做出反应, A 国就有动机在博弈 II 中挑起冲突从而偏离这一策略。

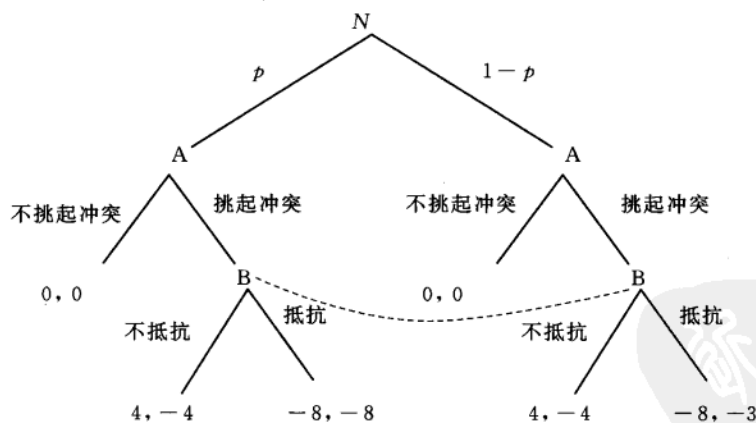


图 8.4 信息结构 III

在本章中, 我们探讨动态环境中信息的策略性使用。不完全信息引发了一些重要问题。

● **信息的策略性使用:** 根据信息分布方式, 有哪一个参与者拥有信息优势? 在许多博弈中, 拥有信息多的一方拥有重大优势。但是, 在有些时候, 信息少的一方

却有优势——无知竟是件好事！

- **学习**: 拥有的信息少的一方通过观察拥有信息多的一方的行动, 能够得到更多信息吗? 这种可能性如何影响拥有信息多的一方所采取的策略呢?

- **信号揭示**: 拥有的信息多的一方能够可信地把关于博弈的信息传递给信息少的一方吗? 拥有的信息多的一方能够误导拥有的信息少的一方吗?

8.1 精炼贝叶斯均衡

子博弈精炼虽能排除掉一些不合理的 Nash 均衡, 但是, 在无法完全地观察到行动的许多扩展式博弈时, 需要更强的均衡概念。我们看图 8.5 中的扩展式博弈。这个博弈给出了雷达发明之前海军舰队的各种部署所带来的收益。参与者 1 选择是否秘密部署军队攻占某个岛屿。参与者 1 可以派遣一支小型舰队(S)或者派遣一支大型舰队(B)或者不派遣任何兵力(ND)。参与者 2 只能观察到对方是否部署兵力。就是说, 借助望远镜, 参与者 2 能够看到对方舰队驶近, 但无法确定舰队规模。如果参与者 1 没有派遣舰队, 参与者 2 就占有这个岛屿, 双方收益是(0, 5); 如果参与者 1 派遣了舰队, 参与者 2 得决定是否迎战。如果不战而降(NR), 参与者 1 占领这个岛屿; 如果参与者 2 抵御对方进攻, 则仍占有此岛屿, 只是在对方派来大型舰队时, 参与者 2 蒙受的损失更高。参与者 1 如果派去的是大型舰队, 自己的损失也会更高。

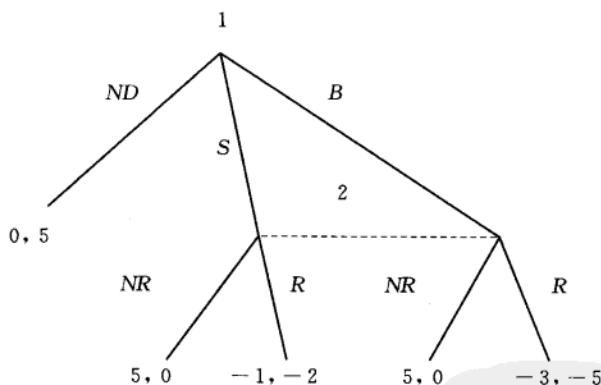


图 8.5 海军部署博弈

这一博弈有三个 Nash 均衡。第一个是(ND, R), 参与者 1 没有出兵, 但是如果出兵, 参与者 2 会抵抗。第二个 Nash 均衡是(B, NR), 参与者 1 派出大舰队, 参与者 2 没有抵抗。另一个均衡是(S, NR)。

这第一个 Nash 均衡却令人莫名其妙。不管参与者 1 派出的是大舰队还是小舰队, 参与者 2 从不抵抗中得到的收益更高。参与者 1 应该看到这一点并因此派出舰队攻岛! 在前一章中, 我们用子博弈精炼剔除了这类问题, 但是在这里, 子博弈精炼概念不起作用。这一博弈没有真子博弈, 所以, (NS, R) 是它的子博弈精炼

Nash 均衡。

这一策略组合之所以不合理,是因为参与者 1 应该预测到参与者 2 在自己的信息集上会做出的理性反应。我们的目的是把这种序贯理性融入均衡概念中。这要求参与者必须对到达每个信息集中的每个历史形成推断,并根据这些推断做出最优反应。这样的均衡被称为精炼贝叶斯均衡(或 PBE)。

回到我们的例子中。对参与者 2 的信息集中的历史形成的任何推断,都不支持 R ,就是说, R 不是最优反应。参与者 2 认为,无论对手使用策略 S 还是使用策略 B ,对手总是派出了舰队。不管是哪种情况,选择 NR 带给参与者 2 的效用都更高。

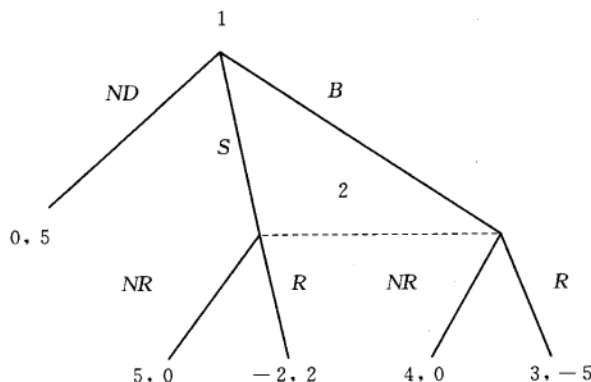


图 8.6 海军部署博弈

这个例子过于简单,我们来看一种更有意思的情况。假设只在选择了 B 时,参与者 1 才能够攻占这个岛屿。而只在参与者 1 选择了策略 S 时,参与者 2 才愿意抵抗对方的进攻。图 8.6 给出了各方的收益。

在这个博弈中, R 或 NR 是否满足序贯理性,决定于参与者 2 对其信息集中的两个可能的历史形成怎样的推断。如果参与者 2 认为对方使用策略 S ,则 R 满足序贯理性;如果参与者 2 认为对方使用策略 B ,则 NR 满足序贯理性。那么,参与者 2 应该形成怎样的推断呢? 参与者 2 的推断显然是以对参与者 1 所采取的行动的预期为基础的。但是,参与者 1 的选择取决于他对参与者 2 形成的推断的预期。

正式定义

信息集可能是单元素集合,也可能是包含多个历史的集合。在扩展式博弈中,参与者 j 被要求在每个信息集 I_j 上采取行动。由于在 h 和 h' 位于同一个信息集时, $p(h) = p(h')$, 所以如果 I_j 是信息集,在这一信息集上采取行动的参与者为 $p(I_j)$ 。给定信息集,我们就能够定义推断。

定义 8.1 在行动无法被完全观察到的扩展式博弈中,信息集 I_j 上的一个推

断是 I_j 上的一个概率分布。一个推断组合是从历史集到单位区间的一个映射 $b: H \rightarrow [0, 1]$, 使得对每个历史 I_j , $b(\cdot)$ 是 I_j 上的推断。对每个 I_j , 如果 I_j 是有限集, 则 $\sum_{h \in I_j} b(h) = 1$; 如果 I_j 是无限集, 则 $\int_{h \in I_j} db(h) = 1$ 。

在前面的例子中, 参与者 2 的信息集上的一个推断是 $\{S, B\}$ 上的一个概率分布。推断组合给出了所有信息集上的推断。由于在每个信息集上只有一个人在行动[即映射 $p(\cdot)$ 是个函数], 所以, 推断组合中的每个推断由谁做出, 是清楚无疑的。在信息集 I_j 上, 如果参与者 i 被要求做出选择, 则推断组合中信息集 I_j 上的推断, 给出了 i 在选择行动时所形成的 I_j 上的推断。

我们现在定义策略需满足的条件, 它被称为序贯理性。不严格地说, 序贯理性要求, 对给定的推断组合, 所有策略在每个信息集上都是最优的。为正式地表述这一思想, 我们用 $p(I_j)$ 表示参与者, 用 $s(I_j)$ 表示信息集 I_j 上的行动。在 $h \in I_j$ 时, 这两个术语等价于 $p(h)$ 和 $s(h)$ 。对给定的策略组合 $s(\cdot)$, 参与者 $p(I_j)$ 在历史 h 上选择 a 而获得的期望效用为 $Eu_{p(h)}(a, h, s(\cdot))$ 。这是个期望效用, 因为其他参与者可能使用混合策略。(在信息集 I_j 上) 如果参与者 $p(I_j)$ 赋予历史 $h \in I_j$ 以概率 $b(h)$, 则在信息集 I_j 上采取行动 a 所实现的期望效用是:

$$E_{p(I_j)}(a, I_j, s(\cdot), b(\cdot)) = \sum_{h \in I_j} b(h) Eu_{p(h)}(a, h, s(\cdot))$$

对给定的推断组合, 如果某个策略组合给出了每个信息集上的最优行动, 则它是序贯理性的。最优性概念要求参与者用 $Eu_{p(I_j)}(a, I_j, s(\cdot), b(\cdot))$ ① 算计行动 a 的合理性。

定义 8.2 在行动无法被完全观察到的扩展式博弈中, 对给定的每个信息集上的推断 $b(\cdot)$, 策略组合 $s(\cdot)$ 在信息集 I_j 上满足序贯理性, 如果对所有行动 s' ,

$$Eu_{p(I_j)}(s(I_j), I_j, s(\cdot), b(\cdot)) \geq Eu_{p(I_j)}(s', I_j, s(\cdot), b(\cdot))$$

如果这个策略组合在每个信息集上都满足序贯理性, 则它满足序贯理性。

回到图 8.6。如果某一推断赋予 S 的概率接近 1, 则在这一信息集上, R 满足序贯理性。如果参与者 2 认为 B 极可能发生, 则 NR 便是序贯理性反应。

要正式地定义精炼贝叶斯均衡, 还需要对推断施加限制条件。推断必须满足弱一致性。就是说, 只要可能, 参与者就应利用贝叶斯法则形成 I_j 上的推断。② 要计算在已到达信息集 I_j 的前提下历史 h 的发生概率, 我们需要知道到达这一信息集的概率。到达 I_j 的概率决定于策略组合。设 $\Pr(I_j | s(\cdot))$ 表示在策略组合 $s(\cdot)$ 下博弈到达 I_j 的概率。如果 h 是 I_j 中的元素, 则在策略组合 $s(\cdot)$ 下, 同时到达 h

① 如果历史集是个无限集, 这里的求和运算得替换为积分运算。

② 前面说过, 贝叶斯法则所产生的是, 在事件 B 发生的前提下, 事件 A 的发生概率。这一条件概率为 $\Pr(A | B) = \frac{\Pr(A \& B)}{\Pr(B)}$ 。

和 I_j 的概率实际上就是到达 h 的概率。我们把这一概率表示为 $\Pr(h|s(\cdot))$, 于是, $\Pr(I_j \& h|s(\cdot)) = \Pr(h|s(\cdot))$ 。因此, 根据贝叶斯法则, 如果在策略组合 $s(\cdot)$ 下博弈到达了信息集 I_j , 则到达历史 $h \in I_j$ 的概率是:

$$\Pr(h|I_j, s(\cdot)) = \frac{\Pr(h|s(\cdot))}{\Pr(I_j|s(\cdot))}$$

定义 8.3 在行动无法被完全观察到的扩展式博弈中, 给定这一博弈的一个策略组合 $s(\cdot)$, 我们说推断 $b(\cdot)$ 相对于策略 $s(\cdot)$ 是弱一致的, 如果在 $\Pr(I_j|s(\cdot)) > 0$ 时, $b(h) = \Pr(h|I_j, s(\cdot))$ 。

将推断的弱一致性与策略的序贯理性结合在一起, 我们就得到了精炼贝叶斯均衡概念。

定义 8.4 在行动无法被完全观察到的扩展式博弈中, 精炼贝叶斯均衡(PBE)是满足以下条件的 $(s(\cdot), b(\cdot))$: (1) 相对于推断 $b(\cdot)$, 策略组合 $s(\cdot)$ 是序贯理性的; (2) 相对于策略组合 $s(\cdot)$, $b(\cdot)$ 是弱一致的。

在 $s(\cdot)$ 下, 如果 I_j 没有被达到, 则 $\Pr(I_j|s(\cdot)) = 0$, 此时, 贝叶斯法则没有定义。弱一致性条件对在这些“偏离均衡”的信息集上形成的推断没有施加任何限制。这一点有时很麻烦; 这意味着可能存在许多不同的精炼贝叶斯均衡, 它们分别以未到达的信息集上形成的不同推断为基础。^①

现在分析图 8.5 中的博弈, 我们要找出 PBE。Nash 均衡 (ND, R) 不是 PBE。给定关于参与者 1 选择 S 还是选择 B 的任何推断, NR 是唯一的序贯理性反应。参与者 2 如果选择 NR , 参与者 1 的最优选择是 S 或 B 。参与者 1 如果选择 B , 则弱一致的推断对参与者 2 在历史 B 上的行动赋予概率 1。于是, 我们得到了一个 PBE: (B, NR) , $\Pr(B|\sim ND) = 1$ 。其中, $\Pr(B|\sim ND)$ 是在参与者 1 没有选择 ND 时, B 的后验概率。使用同样的方法, 我们还能找到另一个 PBE: (S, NR) , $\Pr(S|\sim ND) = 0$ 。

再看图 8.6 中的博弈。参与者 2 如果认为 $\Pr(B|\sim ND) = 1$, 则 NR 是最优反应; 如果认为 $\Pr(B|\sim ND) = 0$, 则 R 是最优反应。一个可能的 PBE 是 (ND, R) 和 $\Pr(B|\sim ND) = 0$ 。由于在参与者 1 选择 ND 时, 弱一致性没有对推断施加任何限制, 所以, 对应着策略 ND , $\Pr(B|\sim ND) = 0$ 是弱一致的。但是, 策略组合 (ND, R) 不满足序贯理性; 参与者 1 如果猜测对手使用策略 R , 就会使用策略 B 而不是 ND 。同样地, ND 也不是对 NR 的最优反应。

① 这里有必要就我们所使用的术语作些说明。对我们所谓的弱一致性, 一些学者使用的术语是一致性; 他们又用这个术语描述序贯均衡中的一个更强的概念。对这两种情况, 我们想用不同术语来表示, 本章后半部分中将定义这个更强的一致性概念。还有一些人在精炼贝叶斯均衡的定义中, 增加了几个偏离路径条件。不正式地说, 这些条件规定: (1) 如果参与者 i 没有关于参与者 j 的更多的信息, 则参与者 i 的行动不会使其他人更新自己关于参与者 j 的类型的推断; (2) 参与者所形成的推断与类型上的同一个联合分布相一致。在只有两个参与者的博弈中, 这些条件并不重要; 如果参与者数量多于两个, 这些条件常常很重要。在政治科学的应用研究中, 极少用到这些条件。因为这个原因, 同时, 也由于在我们给出的 PBE 定义不敷需要时, 我们将使用序贯均衡概念, 所以, 本书没有给出 PBE 的这种定义。感兴趣的读者可以参考 Fudenberg 和 Tirole(1991, 第 8 章)。

假设存在一个纯策略 PBE, 在其中, 参与者 1 没有使用 **ND**。如果参与者 1 使用 **B** 且推断满足弱一致性, 则参与者 2 唯一的序贯理性策略是 **NR**。但是, 参与者 1 如果猜测参与者 2 使用 **NR**, 就会使用策略 **S**。所以, **B** 不可能出现在构成纯策略 PBE 中。参与者 1 如果选择了 **S**, 则弱一致的推断对参与者 2 在历史 **S** 上的行动赋予概率 1。就是说, 唯一满足序贯理性的行动是使用策略 **R**。但是, 参与者 1 如果猜测参与者 2 使用策略 **R**, 参与者 2 就会使用策略 **B**。所以, 在任何纯策略 PBE 中, 都不会用到策略 **S**。因此, 不存在纯策略 PBE。

但是在这一博弈中, 存在混合策略 PBE。假设参与者 1 以概率 q 使用策略 **B**, 以概率 $1-q$ 使用策略 **S**。参与者 2 以概率 z 使用 **R**, 以概率 $1-z$ 使用 **NR**。推断的弱一致性条件要求 $\Pr(B|\neg ND) = q$ 。混合策略均衡则要求在 $\neg ND$ 之后, 参与者 2 在 **R** 和 **NR** 之间无差异, 因此, $q(-5) + (1-q)2 = 0$ 即 $q = \frac{2}{7}$ 。同样的道理, 参与者 1 在 **B** 与 **S** 之间无差异, 因此, $(1-z)5 + z(-2) = (1-z)4 + z3$, 即 $z = \frac{1}{6}$ 。就是说, 给定推断 $\Pr(B|\neg ND) = \frac{2}{7}$, 参与者 1 以 $\frac{5}{7}$ 的概率选择 **S** 和以 $\frac{2}{7}$ 的概率选择 **B**, 与参与者 2 以概率 $\frac{1}{6}$ 选择 **R** 和以概率 $\frac{5}{6}$ 选择 **NR** 的策略组合, 构成 PBE。

8.2 信号揭示博弈

在一类重要的非完美信息博弈中, 拥有较多信息的一方和拥有较少信息的一方在进行博弈。我们称前者为信号发送方, 称后者为信号接收方。在信号发送方首先行动时, 这一博弈被称为信号揭示博弈(signaling game)。这一名称源自于发送方的行动把关于她自身类型的信息传递给接收方。我们用最简单的信号揭示博弈, 来说明博弈参与者所面对的激励, 探讨可能存在的各种均衡。

自然首先为参与者 1 抽取类型 $\theta \in \{a, b\}$ 。参与者 1 观察到自己的类型, 选择信号 $m \in \{a, b\}$ 。参与者 2 能观察到信号但观察不到参与者 1 的类型。在收到信号后, 参与者 2 选择政策 $p \in \{a, b\}$ 。从类型和行动组合中, 每个参与者获得的收益为 $u_i(p, \theta)$ 。图 8.7 中画出了这一博弈的扩展形式。

在这一博弈中, 参与者 2 所面对的信息集中有多个历史。这些集合用虚线表示。参与者 2 在选择政策时, 知道参与者 1 发出的信号, 但是不知道自然所选择的

状态。
在 $u_2(p, \theta) > u_2(p', \theta)$ 和 $u_2(p', \theta') > u_2(p, \theta')$, 其中 $p \neq p'$ 且 $\theta \neq \theta'$ 时, 这一博弈最有意思。不严格地说, 这一条件意味着参与者 2 在选择 p 之前, 想知道 θ 。在双方有相反偏好时, 这一条件得到满足。设 θ 决定每个政策带给每个人的收益: 参与者 1 希望 θ 与 p 匹配, 参与者 2 希望两者不匹配(即 $\theta \neq p$)。为便于分析, 假设参与者 2 是立法机构, 它在现状 a 与新政策 b 之间做出选择。新政策的效果

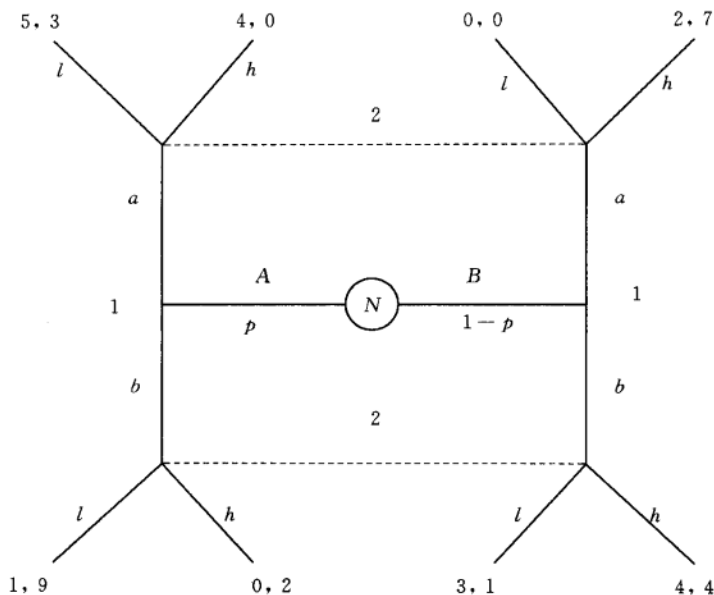


图 8.7 信号揭示博弈

并不确定,但是,立法机构希望新政策导致的结果位于现状点的右方。参与者 1 是拥有信息的专家,专家就 b 所导致的结果是落在现状点左方还是右方,在国会作证,但是,专家的证词是不可证实的。专家并不支持立法机构的偏好;专家所偏好的是所导致的结果落在 a 点左方的政策。 $\theta = a$ 表示的信息为: b 所导致的结果落在 a 点右方。 $\theta = b$ 表示的信息为: b 所导致的结果落在 a 点左方。就是说, θ 是个信号,它告知立法机构实际偏好的是哪一政策。

在这一博弈中,参与者 1 的行动是空谈(cheap talk)^①,因为参与者 2 无法证实参与者 1 发出的信号的准确性。假设说谎没有任何成本。立法机构希望 p 与 θ 相匹配而专家不希望产生这种匹配。这就意味着存在以下收益关系:

$$u_1(b, a) < u_1(a, a)$$

$$u_1(a, b) < u_1(b, b)$$

$$u_2(b, a) > u_2(a, a)$$

$$u_2(a, b) > u_2(b, b)$$

为构造出非完美信息博弈,我们假定状态 θ 上有先验推断。 $\theta = a$ 的概率为 π , $\pi > \frac{1}{2}$ 。

第一个问题是,在这一博弈中是否存在 PBE,参与者 1 即信号发送方,在均衡中向参与者 2 即接收方,揭示自己的私人信息。要存在这一均衡,参与者 1 必须使用以下策略中的一个:

^① 在博弈论和信息经济学中,“cheap talk”指“无成本的信号”,此处译为“空谈”。——译者注

$$m(\theta) = \begin{cases} a, & \text{如果 } \theta = a \\ b, & \text{如果 } \theta = b \end{cases}$$

或者:

$$m(\theta) = \begin{cases} b, & \text{如果 } \theta = a \\ a, & \text{如果 } \theta = b \end{cases}$$

首先探讨其中第一种信号策略。这种形式的策略被称为充分揭示,因为,如果接收方知道这一策略且观察到这一信号,就能推断出发送方的私人信息。这第一种策略也被称为真实的策略。如果参与者 1 使用这一真实的策略组合,推断的弱一致性要求:

$$b(\theta = a | m = a) = \frac{\pi \times 1}{\pi \times 1 + (1 - \pi) \times 0} = 1$$

和

$$b(\theta = a | m = b) = \frac{\pi \times 0}{\pi \times 0 + (1 - \pi) \times 1} = 0$$

给定这些推断,序贯理性要求参与者 2 根据以下映射选择自己的政策:

$$p(m) = \begin{cases} b, & \text{如果 } m = a \\ a, & \text{如果 } m = b \end{cases}$$

要证实上述策略组合和推断构成 PBE,还需要验证前述策略 $m(\cdot)$ 满足序贯理性。其中的关键,是看它是否为对政策映射 $p(\cdot)$ 的最优反应。由于 $u_1(b, a) < u_1(a, a)$ 且 $u_1(a, b) < u_1(b, b)$, 所以,参与者 1 如果观察到 $\theta = a$, 就会偏离上述策略而宣布 $m = b$, 借此导致结果 a ; 如果观察到 $\theta = b$, 就会偏离上述策略而宣布 $m = a$ 。因此,在任何的 PBE 中,都不会使用真实的信号策略 $m(\theta) = \theta$ 。事实上,在任何的 PBE 中,也都不会使用第二种充分揭示的信号策略——这一结论的证明留做习题。前述信号策略被称为分离策略,因为在任何 PBE 中如果发送方使用这种策略,接收方都能够知道 θ 。

这一博弈没有分离均衡,但可能有其他的 PBE。假设无论 θ 是什么,发送者都发送相同的信号,例如 a , 即 $m(\theta) = a$ 。给定这一信号策略,接收方从信号 m 得不到任何信息。在这种情况下,推断的弱一致性要求:

$$b(\theta = a | m = a) = \frac{\pi \times 1}{\pi \times 1 + (1 - \pi) \times 1} = \pi$$

如果信号为 $m = b$, 接收者应该形成怎样的推断呢? 这一问题有些棘手。如果使用贝叶斯法则,就会得到:

$$b(\theta = a | m = b) = \frac{\pi \times 0}{\pi \times 0 + (1 - \pi) \times 0} = \frac{0}{0}$$

$\frac{0}{0}$ 不是个数字,所以贝叶斯法则在这种情况下没有定义。就是说,在表达后验推断

时,作为前提的历史在给定的策略组合中的发生概率为0。弱一致性的定义只要求在分母大于0时,所形成的推断服从贝叶斯法则。由于弱一致性对这里的推断没有施加任何限制,我们可以以有助于找到均衡的任意方式,确定这一后验推断。设 $b(\theta = a | m = b) = \pi$ 。给定这一推断,要确定接收方的哪个策略具有序贯理性,我们只需比较期望效用。如果:

$$\pi u_2(a, a) + (1 - \pi)u_2(a, b) \geq \pi u_2(b, a) + (1 - \pi)u_2(b, b)$$

政策 a 对接收方就更有利;否则,政策 b 更有利。就是说,如果:

$$\pi \geq \frac{u_2(b, b) - u_2(a, b)}{u_2(a, a) - u_2(a, b) - u_2(b, a) + u_2(b, b)}$$

则给定上述推断,在 $p(m) = b$ 时, p 的序贯理性得到满足。如果(8.1)式中的不等关系正好相反,则在 $p(m) = a$ 时, p 的序贯理性得到满足。最后,我们还得验证对给定的政策函数,上述信号策略满足序贯理性。这一点容易验证。由于政策函数在 m 上为常数,所以任何信号函数都是最优反应。这就是说,我们找到了这一信号博弈的均衡。这种形式的均衡被称为混同均衡。

关于 $b(\theta = a | m = b)$ 的假定相当重要。设 π 足够高从而使得接收方的最优反应是 b 。现在,假设 $b(\theta = a | m = b) = 0$ 。序贯理性于是要求 $p(b) = a$ 。就是说,序贯理性要求,在我们改变偏离路径的推断时,偏离路径的政策行动随之改变。

那么,我们能够根据这些新的推断和这一均衡使用混同信号函数 $m(\theta) = a$, 构造出均衡吗? 假设发送方观察到 $\theta = a$ 。在我们所猜想的这一均衡中,发送方发出信号 a 。接收方没有从这一信号中了解到任何信息,他于是选择 b , 因为(在 π 很高时) $\theta = a$ 更可能发生。但是,如果发送方偏离我们所猜想的这一策略而发送信号 b , 接收方的反应是选择政策 a 。我们所猜想的策略和推断组合如果构成 PBE, 发送方就没有动机以这种方式偏离。但是,如果 $\theta = a$, 发送方就会选择偏离以求得到 $p = a$, 而不坚持所猜想的均衡和得到 $p = b$ 。在某些弱一致的推断下,这一混同信号函数可能构成 PBE, 但是却无法在所有的这种推断下都构成 PBE。偏离路径的推断十分重要;它们对在路径上的行为提供了激励。

在这一博弈中,构造混同均衡的另一种方法是允许发送方使用混合策略。假设不管自己的类型如何,发送方都以概率 $\sigma \in (0, 1)$ 发送信息 a 。根据贝叶斯法则,

$$b(\theta = a | m = a) = b(\theta = a | m = b) = \pi$$

和前面一样,接收方的序贯理性策略决定于 π 。在这一使用混合信号的混同均衡中,贝叶斯法则明确了在每个可观察到的信息集上的推断;没有任何信息集位于偏离均衡的路径上。

在应用研究中,学者们常常假设更大的类型空间、信号空间与政策空间。在许多模型中,还可能有多于一个信号发送方和接收方。这些推广增加了模型的复杂性,但是上述概念依然适用。均衡分析的关键是,在给定策略下的弱一致的推断,和验证参与者没有动机偏离均衡策略。在更一般的模型中,除了有分离均衡和混同均衡

外,还可能存在部分分离均衡或部分混同均衡。^①

定义 8.5 在有多个信号发送方和一个信号接收方的一般信号揭示博弈中,设每个发送者有类型空间 Θ_i 和信号空间 M_i ,接收方有 $\Theta = \times_i \Theta_i$ 上的先验推断 $\pi(\cdot)$ 。此时,以下定义适用:在分离均衡中,接收方的后验推断集中在均衡路径上的真实状态上;在混同均衡中,后验推断与均衡路径上的先验推断相一致。在部分分离均衡中,上述两个条件都不成立。

8.3 应用:选举中的阻止进入行为

在对竞选筹款活动的研究中,一个令人不解的问题是,在位的政治家为什么如此卖力地筹措竞选资金,使得资金数量超过其竞选活动所需。标准的解释是,筹资成功的事实反映在位者的竞选实力,在位者通过筹措资金,阻止潜在的竞争对手参选。Epstein 和 Zemsky(1995)对这种解释进行了正式分析。一位挑战者必须决定是否与一位在位的政治家竞选;假设只有当这位在位者在政治上表现脆弱时,挑战者才愿意参选。如果在位者很强,挑战者宁愿不参选。传统观点认为,在位者可能想通过筹措到巨额竞选资金,向挑战者昭示自己的政治实力。如果这笔竞选资金能使挑战者相信在位者实力强大,他就不会参选。

为在模型中反映这一逻辑,我们用 p 表示挑战者(C)认为在位者(I)实力强(S)的先验概率;用 $1-p$ 表示在位者实力弱(W)的概率。为简化分析,假设在当选后,C和I都得到1个单位的效用。假设C战胜W的概率为 π_w ,战胜S的概率为 π_s 。很自然地, $\pi_w > \pi_s$ 。设 k 是C的竞选成本。为使这一模型更有意义,我们假定 $\pi_w > k > \pi_s$,因为如果 $k > \pi_w$,C永远不会参选;如果 $k < \pi_s$,C总是参选。

这一模型的关键假设是,在筹措竞选资金时,S和W发生不同的融资成本。政治实力强的在位者,更容易筹措到资金。为使分析简单,我们只探讨在位者能否筹措到竞选资金;而不考虑在位者筹措到的竞选资金规模。^②就是说,在位者的策略集是 $S_I = \{WC, \sim WC\}$,其中WC是筹措竞选资金,而 $\sim WC$ 是不筹措竞选资金。 $s_I \in S_I$ 表示I的策略。要筹措到竞选资金,类型W和S分别发生成本 c_w 和 c_s , $c_s < c_w$ 。在位者胜出的概率只决定于自己的类型。竞选资金对竞选没有直接影响,但是我们将看到,它可能有间接影响。在观察到I是否筹措竞选资金的决策后,C决定是参加竞选(E)还是不参加竞选($\sim E$)。图8.8用博弈树描述了这一博弈的扩展形式。^③

① 那些看到杯子半空的人,可能喜欢使用部分混同一词;另一些人则可能认为部分分离与自己的生活哲学更一致。我们在这两个术语上无差异。

② Epstein 和 Zemsky(1995)的原始模型放松了这一假定。

③ 我们还可以以类似图8.7的形式表述图8.8中的博弈。我们同时用这两种表述方式,是为了让读者熟悉不同的图形描述方式。

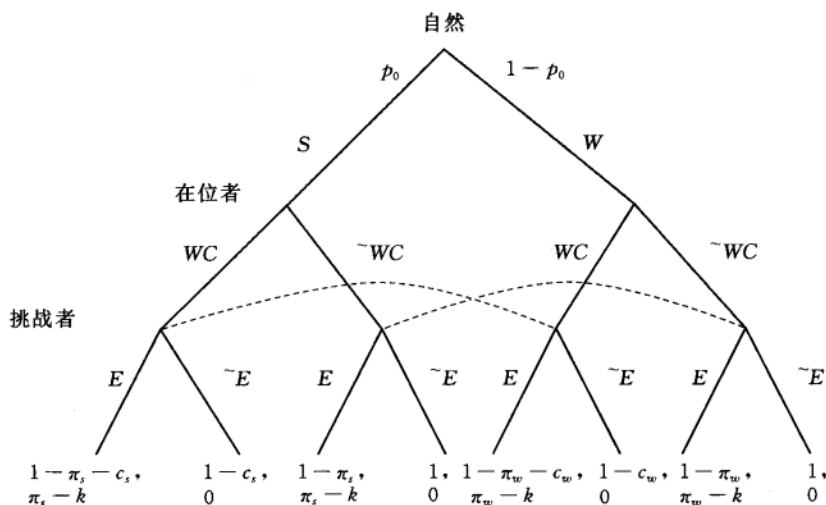


图 8.8 竞选资金博弈

首先看博弈的最后一个阶段。序贯理性要求,只有在进入所带来的期望效用大于或等于 k 时, C 才参选。就是说,参选行为要求 $\pi_s \Pr(S|s_I) + \pi_w \Pr(W|s_I) \geq k$ 。

在第一个阶段, W 和 S 选择是否筹措竞选资金。在这一阶段,有三类可能的均衡:

- (1) 分离均衡: S 和 W 选择不同策略。
- (2) 混同均衡: S 和 W 选择相同策略。
- (3) 半混同策略: S 和 W 选择不同的混合策略。

8.3.1 分离均衡

我们首先分析如下分离均衡:实力强的在位者筹措竞选资金,实力弱的在位者不筹措竞选资金。给定这样的策略组合,挑战者在均衡点上了解到在位者的类型。由于 $\pi_w > k > \pi_s$, 只有当在位者不筹措竞选资金时,挑战者才参与竞选。

命题 8.1 如果 $c_s \leq \pi_s$ 且 $c_w \geq \pi_w$, 以下策略和推断构成精炼贝叶斯均衡: $s_I(S) = WC$, $s_I(W) = \sim WC$, $s_C(\sim WC) = E$, $s_C(WC) = \sim E$, $\Pr\{S|WC\} = 1$, $\Pr\{S|\sim WC\} = 0$ 。

给定这些策略,根据贝叶斯法则, C 的均衡推断是 $\Pr\{S|WC\} = 1$ 和 $\Pr\{W|\sim WC\} = 0$ 。由于 $\pi_w > k > \pi_s$, C 的策略——只在对方没有筹措竞选资金时才进入——是最优反应。我们现在验证在位者的策略也是其最优反应。类型为 S 的在位者从 WC 中得到效用 $1-c_s$ 。由于 $c_s \leq \pi_s$, 所以它大于 $\sim WC$ 所产生的效用 $1-\pi_s$ 。类型为 W 的在位者从筹措竞选资金中得到的效用是 $1-c_w$, 而从不筹措竞选资金中得到的效用是 $1-\pi_w$ 。所以, 给定 $c_w \geq \pi_w$, $\sim WC$ 是最优反应。

分离均衡的存在,为什么要求 $c_s \leq \pi_s$ 且 $c_w \geq \pi_w$ 呢? 竞选资金要能够可信地

揭示出实力强弱,筹措竞选资金的成本,对实力弱的在位者来说,必须比实力强的在位者高出很多。否则,实力弱的在位者也会去筹措资金,于是,通过筹措资金阻止 C 参选的策略不再满足序贯理性。当然了,如果筹措竞选资金的行为无法阻止 C 参选,两种类型的在位者就不会去筹措资金。因此,也就不存在分离均衡。

在这一博弈中,不存在发送相反信号(就是说,实力弱的在位者筹措竞选资金,实力强的在位者不去筹措资金)的均衡。

命题 8.2 如果 $c_s \leq \pi_s$ 且 $c_w \geq \pi_w$, 则策略 $s_I(S) = \sim WC$ 和 $s_I(W) = WC$ 不可能构成精炼贝叶斯均衡。

我们证明这一命题。面对着这样的策略,当且仅当对手使用策略 WC 时, C 的最优反应是参选。于是, S 从 $\sim WC$ 中得到一个单位的效用,从 WC 中得到的效用为 $1 - \pi_s - c_s$ 。所以, S 的策略是最优反应。再看 W 的最优反应。他从 WC 中得到的效用为 $1 - \pi_w - c_w$, 从 $\sim WC$ 中得到一个单位的效用。 $s_I(W) = WC$ 不是最优反应。

不存在这种分离均衡的原因,可以用激励一致性解释。我们说,类型 θ 发送信号 m 而类型 θ' 发送信号 m' 的信号映射,对类型 θ 和 θ' 具有激励一致性,当且仅当 $EU(m|\theta) \geq EU(m'|\theta)$ 且 $EU(m'|\theta') \geq EU(m|\theta')$ 。换句话说,在自己的信号和另一种类型的信号之间,每一种类型都弱偏好自己的信号。在分离均衡中,信号策略必须满足激励一致性。

而如果发送相反的信号的行为要构成均衡,根据激励一致性,就必须有:

$$EU(\sim WC|s) \geq EU(WC|S)$$

$$EU(WC|W) \geq EU(\sim WC|W)$$

这两个不等式相加,移项,得到:

$$EU(\sim WC|s) - EU(\sim WC|W) > EU(WC|s) - EU(WC|W)$$

引入这个模型中的有关概念,上面的条件可以改写为:

$$\Pr\{E|\sim WC\}(\pi_w - \pi_s) > \Pr\{E|WC\}(\pi_w - \pi_s) + c_w - c_s$$

由于 $c_w > c_s$ 且 $\pi_w > \pi_s$, 因此,激励一致性要求 $\Pr\{E|\sim WC\} > \Pr\{E|WC\}$ 。但是,给定 C 的推断,上述策略不可能是最优反应。因此,在 PBE 中,在位者不可能发送相反的信号。在第 11 章机制设计模型中,激励一致性起着核心作用。

8.3.2 混同均衡

现在探讨两种类型的在位者选择相同策略的混同均衡。一般情况下,构造这种均衡的最简单的方法,是在偏离均衡的信号被发送的事件中,找出关于发送方的类型的最不利的推断。在这里的模型中,要找到这种推断,要求 C 在观察到在位者的偏离均衡的行动后,认定在位者是实力弱者。这种悲观的推断如果支持某个策略组合,则这个策略组合构成 PBE。在一些不如此悲观的推断,也可能支持这个策

略组合构成 PBE。作为训练,我们计算支持每个 PBE 的完整的推断集。

在第一个混同均衡中,在均衡路径上,任何类型的在位者都不筹措竞选资金,挑战者不参选。

命题 8.3 设 $p \geq \frac{k - \pi_w}{\pi_s - \pi_w}$ 。以下策略和推断构成精炼贝叶斯均衡: $s_I(W) = s_I(S) = \sim WC$, $s_C(WC) = E$, $s_C(\sim WC) = \sim E$, $\Pr\{S | \sim WC\} = p$, $\Pr\{s | WC\} \leq \frac{\pi_w - k}{\pi_w - \pi_s}$ 。

由于 W 和 S 使用相同的策略,根据贝叶斯法则, $\Pr\{S | \sim WC\} = p$ 。所以,当在位者没有筹措竞选资金时, C 从参选中获得的效用为:

$$\pi_s p + \pi_w (1 - p) - k = \pi_w - (\pi_w - \pi_s) p - k$$

因此,在观察到 $\sim WC$ 后,只有当 $p \leq \frac{k - \pi_w}{\pi_s - \pi_w}$ 时, C 才参选。所以,在 $p \leq \frac{k - \pi_w}{\pi_s - \pi_w}$ 时, $s_C(\sim WC) = \sim E$ 满足序贯理性。如果观察到 WC , C 应该有怎样的推断呢? 他又该怎么做呢? 对偏离均衡路径的历史, PBE 对应该怎么做的问题,没有提供任何答案。为构造出这一具体的 PBE,我们可以有任意的推断,只要它满足 $\Pr\{S | WC\} \leq \frac{\pi_w - k}{\pi_w - \pi_s}$ 。这种推断和序贯理性意味着,在观察到 WC 时, C 参选。再看 S 和 W 的策略。如果不筹措竞选资金,他们都得到 1 个单位的效用;如果筹措竞选资金,他们的效用分别为 $1 - \pi_w - c_s$ 和 $1 - \pi_w - c_w$ 。就是说,他们的策略是最优反应。

另一个混同均衡,在路径上与前一个均衡完全相同,却由不同的偏离路径的推断和策略在支持。

命题 8.4 假设 $p \geq \frac{k - \pi_w}{\pi_s - \pi_w}$ 。以下策略和推断构成精炼贝叶斯均衡: $s_I(W) = s_I(S) = \sim WC$, $s_C(WC) = s_C(\sim WC) = \sim E$, $\Pr\{S | \sim WC\} = p$, $\Pr\{S | WC\} \geq \frac{\pi_w - k}{\pi_w - \pi_s}$ 。

在这里,给定均衡路径上的推断, C 的最优反应是不参加竞选。假设在位者偏离均衡路径而筹措资金。在这一均衡中 C 认为,在偏离均衡的信息集上,在位者实力强的可能性比较大。所以, C 依然选择 $\sim E$ 。由于 C 始终不参选,所以,对两种类型的在位者来说,最优策略是不筹措竞选资金。

在一些混同均衡中,实力弱的在位者通过筹措竞选资金而模仿实力强的在位者。筹措竞选资金的行为阻止 C 参选。

命题 8.5 假设 $p \geq \frac{k - \pi_w}{\pi_s - \pi_w}$, $c_s \leq \pi_s$ 且 $c_w \leq \pi_w$ 。则以下策略和推断构成精炼贝叶斯均衡: $s_I(W) = s_I(S) = WC$, $s_C(\sim WC) = E$, $s_C(WC) = \sim E$, $\Pr\{S | WC\} = p$, $\Pr\{S | \sim WC\} \leq \frac{k - \pi_w}{\pi_s - \pi_w}$ 。

给定上述均衡推断,观察到对手筹措竞选资金, C 不参选显然是最优反应。如果观察到 $\sim WC$, C 认为在位者为类型 S 的概率很低,于是选择 E 。再看类型为 S 的在位者。在均衡中,他得到 $1 - c_s$; 如果偏离均衡,他得到 $1 - \pi_s$ 。因此,只要 $c_s \leq \pi_s$, 类型为 S 的在位者就会选择 WC 。最后看类型为 W 的在位者的选择。在均衡中,他得到 $1 - c_w$; 如果偏离均衡,他得到 $1 - \pi_w$ 。所以,如果 $c_w \leq \pi_w$, WC 是他的最优反应。

还有一类混同均衡。在其中,在位者没有筹措资金,挑战者参选。

命题 8.6 假设 $p \leq \frac{k - \pi_w}{\pi_s - \pi_w}$ 。则以下策略和推断构成精炼贝叶斯均衡: $s_I(W) = s_I(S) = \sim WC$, $s_C(WC) = s_C(\sim WC) = E$, $\Pr\{S | \sim WC\} = p$, $\Pr\{S | WC\} \leq \frac{k - \pi_w}{\pi_s - \pi_w}$ 。

在这里, C 参选是最优反应,因为不管在位者采取什么行动, C 都认为在位者不可能为 S 。在均衡路径上,贝叶斯法则意味着 $\Pr\{S | \sim WC\} = p$ 。在偏离均衡的路径上,我们赋予概率 $\Pr\{S | WC\} \leq \frac{k - \pi_w}{\pi_s - \pi_w}$ 。由于 C 选择参选,所以,任何类型的在位者,如果选择筹措竞选资金,都是在浪费资源。所以,两种类型的在位者都没有动机筹措竞选资金。

许多路径都能构成混同均衡,但是,并不是所有的路径都能具有这一特征。

命题 8.7 在任何均衡中,都不会有 $s_I(W) = s_I(S) = WC$ 和 $s_C(WC) = s_C(\sim WC) = E$ 。

两种类型的在位者都会偏离到 $\sim WC$ 上;这并不会改变 C 的行为,但是却避免了筹措竞选资金的成本。

总之,在这一博弈中,存在两类混同均衡:在一类均衡中,两种类型的在位者都筹措竞选资金;在另一类均衡中,两种类型的在位者都不筹措竞选资金。这两种均衡向挑战者提供相同数量的信息——实际上都没有提供任何信息——但是在在位者所发生的成本上,这两种类型的均衡表现不同。在两种类型的在位者都选择 $\sim WC$ 的均衡和两种类型的在位者都选择 WC 的均衡之间,每个参与者都偏好(弱偏好)前者。就是说,前一种均衡是有效率的PBE,而后一种均衡是无效率的PBE。重要的是,在PBE概念中没有任何内在的性质能使我们预测哪一种均衡更可能发生。但是,正如后面将谈到的,存在着精炼PBE集的各种标准。这些精炼常常能找出有效率的PBE。

8.3.3 部分混同

最后一种可能性是,两种类型的在位者选择不同的混合策略。我们应该探讨这一可能性,原因有几点。首先,如果存在这种均衡,则要全面地描述PBE集,就需要分析部分混同均衡。其次,一些学者有时会谈到信息最充分的均衡(在这种均衡中,信号接收方的后验推断最接近真实的类型分布)。在许多情况中,在一些参

数值上,分离均衡并不存在,但存在部分混同均衡。在这些情况中,信息最充分的均衡便是部分混同均衡。

这里只探讨一种部分混同均衡,而把对其他可能性的讨论留做习题。在这一均衡中,类型为 S 的在位者总是筹措资金,类型为 W 的在位者以概率 q 筹措资金。如果对手筹措资金,则 C 以概率 r 参选。

命题 8.8 设 $\pi_W \geq c_W$ 且 $p \geq \frac{k - \pi_W}{\pi_S - \pi_W}$ 。则以下策略构成精炼贝叶斯均衡:
 $s_I(S) = WC$, $s_I(W) = \{\text{以概率 } q \text{ 选择 } WC\}$, $s_C(\sim WC) = E$, $s_C(WC) = \{\text{以概率 } r \text{ 选择 } E\}$ 。

根据贝叶斯法则, $\Pr\{S | \sim WC\} = 0$, 因此在观察到 $\sim WC$ 后, C 参选是最优反应。如果观察到的是 WC , 根据贝叶斯法则,

$$\begin{aligned}\Pr\{S | WC\} &= \frac{\Pr(WC | S)\Pr(S)}{\Pr(WC | S)\Pr(S) + \Pr(WC | W)\Pr(W)} \\ &= \frac{p}{p + q(1 - p)}\end{aligned}$$

由于在观察到 WC 后, C 使用混合策略,所以在参选和不参选之间他肯定是无差异的。就是说,

$$\pi_S \Pr\{S | WC\} + \pi_W \Pr\{W | WC\} = k$$

代入两个条件概率,得到:

$$\frac{\pi_S p}{p + q(1 - p)} + \frac{\pi_W q(1 - p)}{p + q(1 - p)} = k$$

求解 q , 我们会发现,

$$q^* = \frac{p(k - \pi_S)}{(1 - p)(\pi_W - k)}$$

此时,无差异性条件得到满足。由于 $q^* \leq 1$, 所以,对这一部分混同均衡, $p \geq \frac{k - \pi_W}{\pi_S - \pi_W}$ 是必要条件。 r 同样必须满足无差异条件;在筹措和不筹措竞选资金之间,类型为 W 的在位者无差异。如果选择 $\sim WC$, W 得到 $1 - \pi_W$ 。如果选择 WC , W 的收益为 $r(1 - \pi_W) + (1 - r) - c_W$ 。于是,无差异条件要求:

$$r^* = \frac{\pi_W - c_W}{\pi_W}$$

$r^* \geq 0$ 的前提条件意味着 $\pi_W > c_W$ 。给定这些混合策略,我们只需验证类型为 S 的在位者偏好 WC 。他如果选择 $\sim WC$, 则得到 $1 - \pi_S$; 而在均衡中,他得到 $r(1 - \pi_S) + (1 - r) - c_S$ 。于是, $EU_S(WC) - EU_S(\sim WC) = \pi_S(1 - r) - c_S = \frac{\pi_S c_W}{\pi_W} - c_S$ 。由于

$\pi_s > \pi_w$ 且 $c_w > c_s$, 所以, 这个差大于零。

当 $\frac{c_w}{\pi_w} < 1$ 且 $p \geq \frac{k - \pi_w}{\pi_s - \pi_w}$ 时, 这一部分混同均衡存在。从前面的讨论中我们知道, 只有当 $\frac{c_w}{\pi_w} \geq 1$ 时, 分离均衡才存在。分离均衡和部分混同均衡不能同时存在。但是, 只要 $p \geq \frac{k - \pi_w}{\pi_s - \pi_w}$, 当不存在分离均衡时, 就存在部分混同均衡。

8.4 应用: 信息和立法组织

围绕着立法政治, 一个在不断讨论的话题是立法机构中各种职能委员会的作用。一些学者强调它们在稳定多数通过的投票规则中的作用 (Shepsle, 1979), 一些学者强调它们在维持利益分配性联盟中的作用 (Shepsle and Weingast, 1987; Weingast and Marshall, 1988), 还有一些学者强调它们在促进多数党的利益中的作用 (Cox and McCubbins, 1994)。在 Gilligan 和 Krehbiel (1987) 提出的模型中, 委员会为立法活动提供与政策有关的专业知识。他们的模型与此前所分析的信号揭示博弈有共通之处。

在他们的模型中, 政策制定者并不总是知道政策选择与政策结果之间的准确关系。例如, 立法者可能不知道某一具体的农业政策如何影响农民收入, 因为在气候和外国竞争等因素如何影响农产品价格等问题上, 他们缺少专业知识。立法者重视这种信息, 他们希望有促进信息收集和传递的制度安排。Gilligan 和 Krehbiel 认为, 拥有有限立法权的委员会制度, 有助于解决这类信息问题。

为反映政策选择和政策结果之间的偏差, Gilligan 和 Krehbiel 假定政策结果 x 是政策 p 和随机变量 ω 的可加性函数, 即 $x = p + \omega$ 。这一模型有两个参与者: 议会 F 和委员会 C 。在我们的简化模型中, C 准确地了解 ω , 而 F 的先验分布是:

$$\omega = \begin{cases} \theta, & \text{概率} = 0.5 \\ -\theta, & \text{概率} = 0.5 \end{cases}$$

每个参与者有单维度上的二次方空间偏好: F 有理想点 $f = 0$, C 有理想点 $c > 0$ 。于是, 双方的收益分别为 $u_F(x) = -x^2$ 和 $u_C(x) = -(c - x)^2$ 。 F 并不知道 ω , 所以他从政策 p 中获得的效用为:

$$\begin{aligned} 0.5u_F(p + \theta) + 0.5u_F(p - \theta) &= -0.5(p^2 + 2p\theta + \theta^2) - 0.5(p^2 - 2p\theta + \theta^2) \\ &= -p^2 - \theta^2 \end{aligned}$$

就是说, F 的效用包括两个部分。第一部分是 $-p^2$, 反映了其理想点与期望政策之间的距离差。第二部分是 $-\theta^2$, 反映政策冲击 ω 上的方差。^①由于这一冲击的方

① 在二次方偏好下, 期望效用可以分解为两个部分: 随机变量的期望值带来的效用和随机变量的方差带来的效用。一些人称这一结论为二次方偏好的均值-一方差性质。

差以成本形式进入 F 的效用函数,所以, F 厌恶风险;他愿意“出钱”获得关于 ω 的信息。我们要探讨的是, F 是否会授予 C 超议会权利,借此为其提供专业化——即了解所实现的 ω ——并进而揭示有关信息的动力。正式地说, F 选择处理法案的程序或规则。为简化分析,我们假定 F 在以下两个规则之间做出选择。

(1) 开放规则:委员会提出提案,议会可以自由修改这一提法。

(2) 封闭规则:议会必须在 C 的提案和现状政策之间投票表决。在这里,现状政策为 $SQ = 0$ 。

封闭规则下, C 被授予更大的立法权;他可以对 F 采取“要或者不要”的策略。为把这一问题表述为不完全信息博弈,我们必须明确两种类型——即 $-\theta$ 和 θ ——的委员会的策略,在收到委员会的提案后议会形成的推断,和给定这些推断时议会的最优反应。遵循 Gilligan 和 Krehbiel 的分析方法,我们要找出信息最充分的均衡。就是说,我们将探讨是否存在分离均衡的问题。^①

8.4.1 开放规则

在开放规则下,委员会必须确保它所提供的信息能使提案在议会中通过。假设委员会对每种可能的结果都提出自己的理想政策,借此披露所有的信息——委员会选择 p^c 以使 $x = c$ 或者 $p^c = c - \omega$ 。由于对每种自然状态都提出了不同的政策应对,所以这种提案策略向议会揭示了所有信息。在开放规则下,议会修改提案直至 $x = 0$ 或者 $p^f = -\omega$ 。在这一均衡中,期望效用为 $u_F = 0$ 和 $u_C = -c^2$ 。由于我们想确定是否存在分离均衡,所以,在提交了偏离均衡的草案后,我们要确定对 C 不利的推断。就是说,我们假定 F 对待其他一切提案都好像这些提案来自“高类型”,即 $\omega = \theta$ 。于是,在这种偏离之后,满足序贯理性的策略为 $p^f = -\theta$ 。给定 F 的最优反应,我们能够证明只有三种结果可能出现,那就是 -2θ 、 0 和 2θ 。结果 -2θ 产生自偏离均衡的提案,其时 $\omega = -\theta$ 。结果 0 产生自 $p^c = c - \theta$,其时 $\omega = -\theta$;或者产生自 $p^c = c + \theta$,其时 $\omega = \theta$;或者产生自偏离均衡的任何提案,其时 $\omega = \theta$ 。结果 2θ 产生自 $p^c = c + \theta$,其时 $\omega = \theta$ 。

给定 F 的这些反应,分离是最优反应吗? 在 $\omega = \theta$ 时,如果 C 通过提出 $p^c = c + \theta$ 而选择偏离, F 通过政策 $p^f = \theta$ 。由于这样做导致的政策结果是 2θ ,所以, C 从偏离中得到的效用为 $-(c - 2\theta)^2$ 。在 $\omega = -\theta$ 时,如果 C 通过提出 $p^c = c - \theta$ 而选择偏离, F 通过的政策是 $p^f = -\theta$,导致 $x = -2\theta$ 和 $u_C = -(c + 2\theta)^2$ 。如果 C 提出的是 $p^c \notin \{c + \theta, c - \theta\}$,由于所有这些提案都导致 $p^f = -\theta$,所以, $u_C(\theta) = c^2$ 和 $u_C(-\theta) = -(c + 2\theta)^2$ 。

由于 $c > 0$ 且 $\theta > 0$, 所以,最有吸引力的偏离是表 8.3 中的偏离。

^① 有人可能会说:“哈!你们不是说在分离均衡不存在时,信息最充分的半混同均衡可能存在吗!”有这类疑问的读者应该仔细揣摩习题 8.7。这里的情况不属于前述那些情况。

表 8.3 可能的偏离

偏 离	条 件	效 用
$p^c = c + \theta$	$\omega = \theta$ 且 $c > \theta$	$-(c - 2\theta)^2$
$p^c \notin \{c + \theta, c - \theta\}$	$\omega = \theta$ 且 $c > \theta$	$-c^2$
$p^c \neq c + \theta$	$\omega = -\theta$	$-(c + 2\theta)^2$

“低”类型的 C 不会选择偏离, 因为 $-c^2 > -(c + 2\theta)^2$ 。如果 $c < \theta$, “高”类型的 C 也不会选择偏离; 其均衡效用和他在最有利的偏离上得到的效用相同。但是, 如果 $c > \theta$ 且 $\omega = \theta$, 相对于自己的均衡收益, “高”类型的 C 偏好 $-(c - 2\theta)^2$ 。因此, 在 $c > \theta$ 时, 不存在分离均衡。没有分离均衡, 唯一的均衡就是不能揭示任何信息的混同均衡——两种类型的委员会在提案集上使用相同的混合策略。这就是说, 在 $c > \theta$ 时, 没有任何信息被揭示出来。

8.4.2 封闭规则

在开放规则下, 信息无法被委员会揭示出来, 因为议会利用这一信息左右委员会, 并推动政策靠拢议会中间派的理想点。但是, 在封闭规则下, 委员会就不会如此被动。议会或者接受提案, 或者拒绝提案而支持现行政策。假设存在分离均衡, 在均衡中, F 了解到 ω 的值。议会从现行政策中得到的效用是 $-\omega^2$, 从任意提案 p 中得到的效用是 $-(p + \omega)^2$ 。这意味着 F 接受 0 到 -2ω 区间中的任何提案。所以, 在 $\omega = \theta$ 时, F 所接受的最大的政策是 0, 从而导致政策结果 θ 。在 $\omega = -\theta$ 时, F 所接受的最大政策是 2θ , 这也导致政策结果 θ 。

所以, 在 $\omega = \theta$ 时通过提案 $p^c = 0$, 在 $\omega = -\theta$ 时通过提案 $p^c = 2\theta$, 委员会能确保得到结果 θ 。如果 $c < \theta$, 在 $\omega = \theta$ 时通过提案 $p^c = c - \theta$, 在 $\omega = -\theta$ 时通过提案 $p^c = c + \theta$, 委员会能确保得到结果 c 。

要完整描述这一精炼贝叶斯均衡, 我们还必须明确在 C 提出偏离均衡的提案后, F 的行为和推断。要支持所有可能的分离均衡, F 必须接受满足条件 $-2\theta \leq p^c \leq 0$ 的所有提案并否决任何提案 $p^c > 0$ 。这一策略是对在偏离均衡的提案之后形成的推断 $\Pr(\omega = \theta | p^c) = 1$ 的最优反应。

现在验证 C 偏好均衡策略。在 $c < \theta$ 时, C 得到其理想点, 因此无法通过偏离做得更好。于是, 和开放规则下一样, 存在分离均衡。所以, 我们只需探讨 $c > \theta$ 时的情况。设 $\omega = -\theta$ 。此时, 从提案 $p^c = 2\theta$ 中, C 得到 $-(c - \theta)^2$; 从其他提案中, C 得到 $-(c + \theta)^2$ 。在这种情况下, C 不会选择偏离。设 $\omega = \theta$ 。此时, 从提案 $p^c = 0$ 中, C 得到 $-(c - \theta)^2$; 从提案 $p^c = 2\theta$ 中, C 得到 $-(c - 3\theta)^2$; 从其他提案中, C 得到 $-(c + \theta)^2$ 。所以, 只要 $c < 2\theta$, 委员会弱偏好其均衡提案 $p^c = 0$ 。如果 $c > 2\theta$, 委员会就会选择偏离。在这种情况下, 唯一的均衡没有揭示出任何信息; 两种类型的委员会在所有的提案上使用相同的混合策略, 议会则否决 $p^c = 0$ 以外的所有提案。

在开放规则无法实现分离均衡时(即在 $2\theta > c > \theta$ 时), 封闭规则能够实现分离均衡。因此, 根据这一模型, 限制性规则能够鼓励从委员会到议会的更高层次的信息传递。

8.4.3 委员会专业化

在 Gilligan-Krehbiel 模型中, 相对于开放规则, 限制性规则能实现从拥有信息的委员会到议会的更高层次的信息传递。我们现在分析, 议会坚守限制性规则, 能否诱使委员会专业化。现在的博弈采取以下形式:

(1) F 决定委员会是在开放规则下还是在封闭规则下提交提案。

(2) C 决定是否通过发生成本 k 以了解 ω 而实现专业化。

(3) C 提出政策 p^c 。

(4) F 观察到委员会的专业化决策和提案 p^c , 做出政策选择。在封闭规则下, F 在政策 p^c 和现行政策 SQ 之间投票做出选择。在开放规则下, F 选择 p^c 。

如果委员会没有专业化, F 根据自己的先验推断做出决策。因此, 在封闭规则下, 他否决除 $p^c = 0$ 以外的任何政策。在开放规则下, F 通过 $p^f = 0$ 。就是说, 在这两种规则下, 委员会的无专业化选择都导致政策 $p = 0$ 。

现在分析开放规则下委员会的专业化决策。在 $c \leq \theta$ 时, 如果委员会专业化, 则 C 和 F 使用开放规则下的分离均衡策略, 委员会的收益是 $-c^2 - k$ 。如果委员会没有专业化, 它从政策 $p = 0$ 中得到的期望效用是 $-c^2 - \theta^2$ 。所以在 $k \leq \theta^2$ 时, 委员会专业化。如果 $c > \theta$ 且委员会专业化, 则委员会和议会使用混同均衡策略, 导致政策 $p = 0$, 给委员会带来的收益为 $-c^2 - \theta^2 - k$ 。就是说, 在这些条件下, 委员会没有专业化。

再看封闭规则。如果 $c \leq \theta$, 分离均衡产生结果 c , 因此, 委员会从专业化中得到的效用为 $-k$ 。如果 $c > \theta$, 专业化带给委员会的效用为 $-(c - \theta)^2 - k$ 。在这两种情况中, 非专业化导致的收益都为 $-c^2 - \theta^2$ 。通过比较这些效用, 我们能够看出, 在 $k \leq 2\theta c$ 和 $c > \theta$ 时或者在 $k < c^2 + \theta^2$ 和 $c \leq \theta$ 时, 委员会应该专业化。由于在这两种情况中 k 的临界值都高于 θ^2 , 所以, 在开放规则下委员会不会选择专业化, 在封闭规则下委员会常常专业化。

8.4.4 结论

通过计算不同规则下议会得到的效用, Gilligan 和 Krehbiel 得到了关于制度设计的几个结论。在开放规则下, 议会的期望收益是:

$$\text{如果 } c \leq \theta \text{ 且 } k \leq \theta^2, E[u_F] = 0$$

$$\text{在其他条件下, } E[u_F] = -\theta^2$$

在封闭规则下, 议会的期望收益是:

如果 $c \leq \theta$ 且 $k \leq c^2 + \theta^2$, 则 $E[u_F] = -c^2$

在其他条件下, $E[u_F] = -\theta^2$

这里有几点特征值得指出。首先, 当 c 的理想点趋近议会的理想点时, F 的效用总是增加。这是 Krehbiel(1991) 的观点的基础。他认为, 多数党控制的立法机构应该任命能代表议会偏好的委员会。Krehbiel 的结论与关于立法机构的利益分配理论解释不同, 这种解释认为, 委员会由那些对该委员会辖权内的政策有高需求的人构成。

其次, 这一模型指出了 F 偏好封闭规则的前提条件。具体地说, 当 $c \leq \theta$ 且 $\theta^2 \leq k \leq c^2 + \theta^2$ 时, F 选择封闭规则。就是说, 其偏好类似议会中间派偏好的委员会, 和专业化成本适中的委员会, 最有可能被给予封闭规则。

8.5 应用: 信息沟通目的的游说活动

关于利益集团政治的文献发现了一个令人惊讶的经验规律: 利益集团几乎总是游说他们的朋友。由于不管有无游说活动, 这些朋友都投票支持利益集团, 所以这一发现常常被理解为, 游说活动不是政策过程的重要组成部分。

在 Austen-Smith 和 Wright(1992, 1994) 的模型中, 利益集团确实游说他们的朋友。但是, 在均衡点上, 这种游说很重要。这一模型的主要前提是, 利益集团拥有关于立法决策的结果的私人信息。立法者缺乏这方面的信息。在这一模型中, 游说活动表现为利益集团向立法者陈述立场。由于游说活动要发生成本, 所以, 利益集团选择是否游说立法者。Austen-Smith 和 Wright 的主要结论是, 利益集团常常游说友好的立法者, 目的是为了消除对手的游说活动的影响。

在介绍完整模型之前, 我们首先探讨只有一个利益集团和一个立法者的简单模型。立法者必须在政策 A 和 B 之间做出选择。但是, 她不知道自己偏好哪一政策。假设存在两种自然状态 $s \in \{A, B\}$ 。在状态 A 中立法者偏好政策 A , 在状态 B 中偏好政策 B 。为使分析简单, 我们假定在状态 A 中, $u_L(A) = 1$; 在状态 B 中, $u_L(A) = -1$ 。在两种状态中, $u_L(B) = 0$ 。在先验概率分布中, $s = A$ 的概率为 $p < \frac{1}{2}$ 。在这种情况下, 如果没有游说活动所提供的信息, 立法者就选择政策 B 。

在政策 A 和政策 B 之间, 利益集团 G_A 偏好政策 A 。他在两种状态中的效用分别是 $u_{G_A}(A) = 1$ 和 $u_{G_A}(B) = -1$ 。如果 G_A 决定游说, 他得付出成本 $c > 0$ 才能准确了解到真实状态。他的先验推断与立法者相同。一旦获得了信息, 利益集团就会发出信号 $m = A$ 或 $m = B$ 。这里的信号可以任意解释。

在观察到信号 m 后, 立法者 L 可以通过审查所收到的信号而验证利益集团的信息。如果这样做, L 就得发生成本 κ 。如果信号被发现不正确(就是说, $m \neq s$), 利益集团就会被罚款 δ 元。

现在看利益集团的可能的游说策略。首先假设在任何状态 s 中, G_A 总是报告 A 。就是说, 这一信号没有提供任何信息。如果不审查信号, L 就会始终选择政策 B 。

如果他审查信号,在审查之后, L 选择最优结果。由于从审查活动中获得的效用是 $1-\kappa$,所以,当且仅当 $p > \kappa$,在 L 对利益集团“总是报告 A ”的策略的最优反应中,包含信号审查。

由于“总是报告 A ”的游说策略不影响结果,所以, G_A 宁愿选择不游说。因此,成功的游说策略要求 G_A 至少有时候报告真实状态。如果 G_A 总是报告真实状态, L 就会始终对政策 A 言听计从并且永远不进行信号审查。这就使得 G_A 有动机,甚至在 $s=B$ 时,报告状态为政策 A 。由于 G_A 有动机偏离,所以这一策略不是均衡策略。

所以,这一博弈的任何精炼贝叶斯均衡都是部分混同均衡。设在某个均衡中,在 $s=A$ 时, G_A 报告状态为 A ;在 $s=B$ 时, G_A 以 μ 的概率报告状态为 A 。给定这一策略,在收到 A 的信号后, L 利用贝叶斯法则更新自己对 $s=A$ 的概率推断为:

$$\hat{p} = \frac{p}{p + \mu(1-p)}$$

在这一均衡中, L 在投票支持政策 A 和对信号进行审查之间无差异。投票支持政策 A 实现的期望效用为 $\hat{p} - (1 - \hat{p}) = 2\hat{p} - 1$,进行审查实现的期望效用为 $\hat{p} - \kappa$ 。所以, μ 必须满足条件:

$$\frac{p}{p + \mu(1-p)} = 1 - \kappa$$

解得 μ 为:

$$\mu = \frac{\kappa p}{(1-p)(1-\kappa)}$$

设 L 以概率 α 进行信号审查,从而使得在 $s=B$ 时, G_A 在说谎和说真话之间无差异。 G_A 如果说真话,得到的效用为 -1 ; G_A 如果说谎,得到的效用为 $-\alpha(\delta+1) + (1-\alpha)$ 。使这两个期望效用相等的 α 值是 $\alpha = \frac{1}{\delta+2}$ 。

给定这些游说和信号审查策略,我们计算 G_A 的期望效用,以确定它是否寻找信息和游说。在没有游说时, L 总是选择政策 B ,所以,利益集团从不游说的策略中获得的效用是 -1 。如果他游说,在 $s=A$ 时,他肯定得到政策 A ;在 $s=B$ 时,他得到期望效用 -1 。因此,他从游说活动中得到的事前期望收益为 $2p-1-c$ 。当且仅当 $p > \frac{c}{2}$ 时,利益集团选择游说。

最后验证这样的混合策略概率是否确实是概率。要使 $0 \leq \mu \leq 1$,要求 $1-p > \kappa > 0$;就是说,对立法者而言,信号审查成本不能太高。如果审查成本很高,立法者就不会进行审查,而利益集团就有动机说谎。但是,如果利益集团总是说谎,游说活动就没有效果,利益集团于是不会去寻找信息。

现在看包含两个利益集团的模型。利益集团 G_B 的偏好与 G_A 正好相反: $u_{G_B}(B) = 1, u_{G_B}(A) = -1$ 。要想知道真实状态,它也得付出成本 c 。在第一个阶段中,两个集团同时决定是否游说。一个集团如果选择游说,它就得发生成本 c 以获

取信息。在第二个阶段,前一阶段中选择游说的团体观察到对手的决定,选择自己的信号。两个团体如果都游说,则信号 m_A 和 m_B 同时发出。

由于在没有游说时 L 选择政策 B ,所以不存在 G_B 游说而 G_A 不游说的均衡。如果只有 G_B 游说,它发生成本 c 但却无法改变这一结果。Austen-Smith 和 Wright 对此的解释是,利益集团只游说“友好的”立法者,以对抗对手的游说活动。

前面阐述了只有 G_A 游说时的结果,这里则只探讨两个团体都游说时的结果。我们来看下面的均衡。两个团体都发送真实信号。当 $m_A = m_B = s$ 时,立法者认为真实状态是 s 。如果 $m_A \neq m_B$, L 就对信号进行审查并选择自己的最优政策。为使信号审查决策构成对偏离均衡的信号的最优反应,我们假定 $\Pr(s = A | m_A \neq m_B) \geq \kappa$ 。

我们首先证明, G_A 不会偏离这一均衡,即不会发送虚假信号。在 $s = A$ 时,他显然没有动机选择 $m_A = B$ 。我们分析 $s = B$ 时,他是否会选择 $m_A = A$ 。发送信号 $m_A = A$, 会招致信号审查,从而导致政策 B 和罚款 δ , 所以,这样做带给他的总收益是 $-1 - \delta$ 。而说真话实现的收益是 -1 , 所以, G_A 没有动机偏离均衡策略。

再看 G_B 的决策。我们需要证明在 $s = A$ 时,他不会选择 $m_B = B$ 。和上一段落中的分析一样,对 G_B 来说,发送信号 $m_B = B$ 会招致信号审查和罚款。所以,他不会偏离均衡策略。

再来分析在什么条件下,两个利益集团都游说。两个集团如果都游说,就会导致信息充分的结果,所以 G_A 和 G_B 的均衡效用分别是 $2p - 1 - c$ 和 $1 - 2p - c$ 。如果只有 G_A 游说, G_B 得到的效用为:

$$-p + (1-p)[\mu(2\alpha - 1) + (1-\mu)] = 1 - 2p - \left[\frac{\kappa p}{1-\kappa} \frac{2\delta + 2}{\delta + 2} \right]$$

如果:

$$p > \frac{1-\kappa}{\kappa} \frac{\delta + 2}{2\delta + 2} c$$

这一效用就小于 $1 - 2p - c$ 。前面说过,如果只有 G_A 在游说,就没有揭示任何信息,立法者选择政策 B 。这给 G_A 带来的效用为 -1 。所以,只要 $2p - 1 - c \geq -1$, G_A 就会游说。这一条件正是 $p \geq 2c$, 即此前得到的条件。因此,无论 G_B 是否游说, G_A 都会选择游说。

我们对两个利益集团的均衡游说决策概括如下:

- (1) 如果 $p < \frac{c}{2}$, 两个集团都不游说;
- (2) 如果 $\frac{1-\kappa}{\kappa} \frac{\delta + 2}{2\delta + 2} c > p > \frac{c}{2}$, 只有 G_A 游说;
- (3) 如果 $p > \frac{1-\kappa}{\kappa} \frac{\delta + 2}{2\delta + 2} c$, 两个集团都游说。

Austen-Smith 和 Wright 将这个精炼贝叶斯均衡概括为:

- 在其他条件相同时,当立法者就某个议题收到来自某一方的团体游说时,进行游说的团体是反对立法者事前政策立场的团体。

- 一个集团的游说“不友好的”立法者的决策,独立于对立团体的游说决策。
- 一个利益集团,在观察到与自己“友好的”立法者在被对立团体游说时,所做出的游说此立法者的决策,纯粹是为了抵消对手的影响。

8.6 对精炼贝叶斯均衡的提炼*

8.6.1 序贯均衡

从前面几个例子可以看出,弱一致性概念可谓名实相符。对偏离路径的推断,它常常没有施加足够强的限制。实际上,PBE 甚至未必是子博弈精炼均衡,图 8.9 中的博弈说明了这一点。

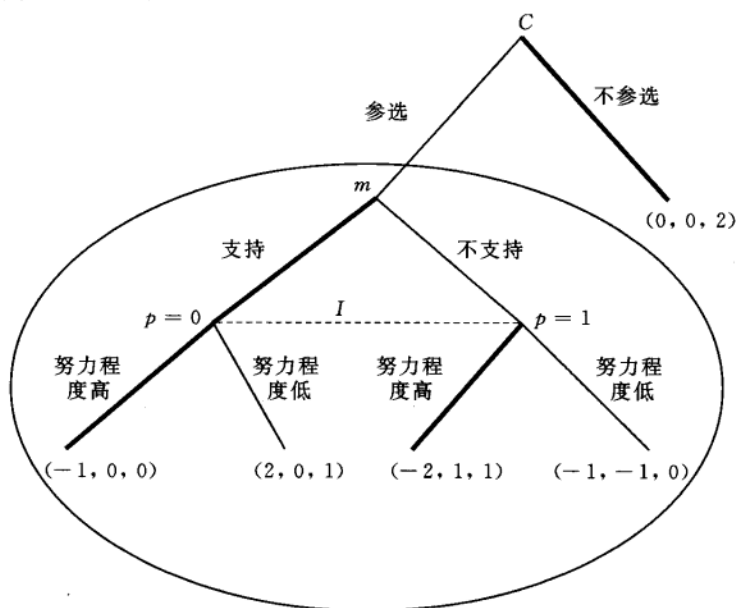


图 8.9 竞选博弈

在这个扩展式博弈中, C 是某一公职的潜在竞选人,他要决定是否参选。在做出参选决策后,当地媒体决定是否支持 C 。在媒体做出支持决策之后但在观察到这一决策之前, I 必须决定投入多大精力到竞选活动中。^①根据图中的收益数字,只在得到媒体支持且在位者投入的努力程度低时, C 才会参选;只在在位者选择的努力程度低时,媒体才支持 C ;只在 C 没有得到媒体支持时,在位者才愿意投入精力到竞选中。

在图 8.9 中,粗线路径为一个策略组合, $p=0$ 和 $p=1$ 为推断。不难看出,这一策略组合和推断构成一个 PBE。根据给定的推断,媒体不支持 C 的概率为 1,所

^① 也可以认为媒体的支持决策和在位者的努力决策同时做出。

以,在位者最优地选择高努力水平。而且,预期到在做出参选决策后,接踵而来的是媒体的支持和在位者的高投入,因此对 C 来说,最优决策是不参选。在这一策略组合下,在位者的信息集无法被到达,且弱一致性没有对所做出的推断施加任何限制。尽管这是个 PBE,但是其中的策略组合并不构成子博弈精炼。为说明这一点,我们来看图 8.9 中圆框内自媒体 m 的决策开始的子博弈。由于 m 支持 C ,所以,在位者的高投入不是最优反应。这里的问题是,对某些策略组合,存在着偏离均衡路径的多节点信息集。

在一个合理的均衡中, I 关于 m 的决策的推断(当然了,前提是如果博弈到达了 I 的信息集)应该与参与者 m 的策略相一致。我们可以猜测:(1)如果 C 参选,则或者 m 支持他且 I 投入的努力程度低;或者(2) m 不支持他且 I 投入的努力程度高;或者(3) m 和 I 随机做出选择。将序贯理性与更强的一致性概念相结合,就能防止出现上述问题。

在本节中,我们只在有限博弈背景中,给出几个更强的均衡概念。所有这些概念都包含定义 8.2 中所定义的序贯理性。它们之间的差异,源于施加于推断上的限制条件。

我们首先引入一些符号。在有限扩展式博弈中,混合策略组合 $\sigma(\cdot)$ 是个映射,它给出了决定着每个信息集所对应的各种可能的行动上的概率分布。对有限博弈,这种策略组合可以表示为向量 $(\sigma(1, 1), \dots, \sigma(a, I), \dots)$, 坐标 $\sigma(a, I)$ 表示在信息集 I 上使用行动 a 的概率。一个策略组合,如果在每个信息集上以正概率选择每个行动,就是完全混合的策略组合。一个纯策略于是就是只由 0 和 1 构成的向量。我们说,混合策略列 $\{\sigma(\cdot)^n\}_{n=1}^{\infty}$ 收敛于混合策略 $\sigma(\cdot)$, 如果对每个 I 和 a , $\sigma(a, I)^n$ 收敛于 $\sigma(a, I)$ 。序贯均衡概念使用比弱一致性概念更强的限制条件。

定义 8.6 在行动无法被完全观察到的有限扩展式博弈中,序贯均衡(SE)是 $(\sigma(\cdot), b(\cdot))$, (1)给定推断 $b(\cdot)$, 策略组合 $\sigma(\cdot)$ 满足序贯理性;(2)存在某个完全混合策略列 $\{\sigma(\cdot)^n\}_{n=1}^{\infty}$ 和推断列 $\{b(\cdot)^n\}_{n=1}^{\infty}$ 分别收敛于 $\sigma(\cdot)$ 和 $b(\cdot)$, 并且存在某个 n' , 如果 $n > n'$, 则给定策略组合 $\sigma(\cdot)^n$, $b(\cdot)^n$ 是弱一致的。

现在说明这一概念是如何精炼我们的 PBE 概念的。回到图 8.9 中的博弈。给定任意的完全混合策略组合列 $\{\sigma(\cdot)^n\}_{n=1}^{\infty}$, 设在 $\sigma(\cdot)^n$ 中, 每个纯策略被以 ε_n 的概率在使用。设在 $n \rightarrow \infty$ 时, $\varepsilon_n \rightarrow 0$ 。如果这一混合策略列收敛于图 8.9 中的 PBE 上, 则 σ^n (支持, 参选)肯定是对 σ^n (参选) $= \varepsilon_n$ 和 σ^n (高努力, 参选) $= 1 - \varepsilon_n$ 的混合最优反应。媒体从支持行为中得到的期望效用因此是 ε_n , 从不支持行为中得到的期望效用为 $1 - 2\varepsilon_n$ 。因此, 随着 ε_n 收敛于 0, 肯定存在某个 n' 使得对所有的 $n > n'$, m 选择 σ^n (支持, 参选) $= \varepsilon_n \rightarrow 0$ 。图 8.9 中的均衡并不具有这一特征, 所以它不是完全混合策略的极限, 也就不是序贯均衡。

序贯均衡还对推断施加了限制。弱一致性要求, 关于媒体的决策的推断与这一策略相一致, 即 b^n (支持) $= \sigma^n$ (支持, 参选)。因此, 给定一个完全混合的策略列, 如果它收敛于以 1 的概率使用策略(不竞选, 如果竞选就支持)的策略组合, 则

$b(\cdot)$ 必须收敛于赋予支持策略的概率为 1 的概率推断上。因此,没有完全混合的策略列和弱一致的推断列,收敛于图 8.9 中给出的策略和推断上。事实上我们能够证明,每个序贯均衡都包含子博弈精炼策略。作为习题,请你证明这一结论。

序贯均衡是对 PBE 的改进,但在许多情况中,这两个概念产生相同的均衡集。例如,在前面讨论的信号揭示博弈中,这两个概念完全一致。Fudenberg 和 Tirole (1991)证明了下面的等价性结论。

命题 8.9 (Fudenberg and Tirole, 1991) 一个有限的扩展式博弈如果只包含两个阶段或者每个参与者至多有两种类型,则 PBE 集和 SE 集相同。

在有限博弈中,始终存在序贯均衡。这是序贯均衡的一个相当好的性质。并且,这一相当好的性质本身也有一个相当好的性质——序贯均衡的存在性证明,同时阐述了标准式博弈中的精炼均衡与扩展式博弈中的序贯均衡之间的关系。

前面说过,我们总能把扩展式博弈表述为标准式博弈:参与者的一个策略规定了他在自己的每个信息集上的行动;每个参与者同时选择自己的策略。第 7 章开始时的例子说明了这种等价性。就我们的目的而言,我们需要探讨另一种略为不同的标准式表述。

定义 8.7 给定一个扩展式博弈,它的代理人标准式表述形式,赋予在各个信息集上行动的参与者不同的身份。

例如,在图 7.2 的蜈蚣扩展式博弈中,代理人标准式表述形式就涉及五个参与者。设 $\sigma(\cdot)$ 是代理人标准式表述中的任意策略组合。在这里,我们用 $\sigma(\cdot)$ 同时表示扩展式博弈中相对应的策略组合。如果参与者 i 在扩展式博弈的多个信息集上行动,则其策略由代理人标准式表述中的多个参与者的策略构成。

把扩展式博弈转化为标准形式,我们就能对均衡集进行如下比较。

命题 8.10 在一个有限扩展式博弈中,如果 $\sigma_a(\cdot)$ 是这一博弈的代理人标准式表述中的精炼均衡,则存在推断 $b(\cdot)$ 使得 $(\sigma(\cdot), b(\cdot))$ 构成扩展式博弈中的序贯均衡。

这一结论除了提供这两个概念之间的联系外,还有一个原因使得它很重要。Selten(1965)证明了有限博弈拥有精炼均衡,所以,我们能直接得到下面结论。

命题 8.11 (Selten) 每个有限扩展式博弈有至少一个序贯均衡。

8.6.2 直觉标准

序贯均衡的概念有时能剔除一些不可行的 PBE,但是,未必总能做到这一点。所以,Kreps 和 Cho(1987)提出了直觉标准,进一步对某些信号揭示博弈中的均衡集施加限制。Kreps 和 Cho 指出,推断应该集中在最有动机偏离的信号发送者身上。他们用 Beer & Quiche 博弈阐述他们的思想。由于这个博弈可能是这方面的最简单的例子,所以,这里也使用这个例子,但加入一些政治因素。设想在美国进攻伊拉克前夕,萨达姆·侯赛因和乔治·布什之间的博弈。这是个两人信号揭示

Figure 1 is a game tree diagram illustrating a signaling game between Bush and Saddam. The game starts with a chance node N , which branches into two types of Saddam: $p = \frac{1}{4}$ and $1-p = \frac{3}{4}$. Saddam then chooses between a and b . Bush then chooses between y and n . Payoffs are given at the end of each branch.

The payoffs for Bush are:

- Top-left: $0, 1$
- Top-right: $1, 0$
- Bottom-left: $1, 0$
- Bottom-right: $2, 1$

The payoffs for Saddam are:

- Top-left: $2, 0$
- Top-right: $3, 1$
- Bottom-left: $3, 0$
- Bottom-right: $2, 1$

The game tree is divided into two regions by a dashed line, labeled "布什" (Bush) at the top and "萨达姆" (Saddam) at the bottom. The nodes are labeled a , b , y , and n . The chance node N is in the center, with branches labeled w and $\sim w$.

图 8.10 布什—萨达姆博弈

在这一博弈中,存在一些混同均衡。在一些混同均衡中,萨达姆允许武器核查,即 $s_h = y$; 在另一些混同均衡中,萨达姆禁止武器核查,即 $s_h = \sim y$ 。在任何混同均衡中,布什的后验推断必须与其前验推断相一致。因此,如果 $\theta = w$ 的先验概率很低,布什就不进攻。 $s_h = y$ 要成为混同均衡策略,两种类型的萨达姆都必须偏好 y 而不是 $\sim y$ 。这就有了激励一致性条件:

$$3 \geq EU_H(\sim y, \sim w)$$

$$2 \geq EU_H(\sim y, w)$$

设 $\Pr(a|\sim y)$ 表示在观察到 y 后布什发起进攻的概率。上述激励一致性条件可表示为：

$$3 \geq 0\Pr(a|\sim y) + 2(1 - \Pr(a|\sim y))$$

$$2 \geq 1\Pr(a|\sim y) + 3(1 - \Pr(a|\sim y))$$

只要 $\Pr(a|\sim y) \geq \frac{1}{2}$ ，这些条件就成立。如果在布什使用的策略中， $\Pr(a|\sim y) \geq \frac{1}{2}$ ，则在出现偏离路径的行动 $\sim y$ 时，他所形成的关于萨达姆类型的后验推断必须满足 $\Pr(w|\sim y) \geq \frac{1}{2}$ ，但是，这一后验推断无法用贝叶斯法则得到。路径上的推断 $\Pr(w|y) = \frac{1}{4}$ 则与先验推断相同，并在混同均衡中成立。所以，我们证明了存在 PBE(根据命题 8.9，这意味着存在 SE)，在均衡中，两种类型的萨达姆使用同一策略 y ，而布什选择不进攻。

两种类型的萨达姆都选择 $\sim y$ 且布什选择不进攻的策略要构成混同均衡，只要求偏离路径的推断 $\Pr(w|y) \geq \frac{1}{2}$ 。给定这种推断，布什以至少 $\frac{1}{2}$ 的概率发起进攻，所以两种类型的萨达姆都会选择 $\sim y$ 借此避免布什进攻。

这两类均衡虽然都是 PBE 和 SE，但是 Kreps 和 Cho 指出，其中只有一类混同均衡是合理的。我们分析两种类型的萨达姆都选择 $\sim y$ 的均衡。这一均衡要求偏离路径的推断满足 $\Pr(w|y) \geq \frac{1}{2}$ 。如果萨达姆允许武器核查，布什会认为他更可能拥有大规模杀伤性武器吗？在这一均衡中，类型 w 得到最大收益。如果没有武器核查也没有进攻，他的收益为 3。如果他偏离到 y 且布什没有进攻，他的收益是 2。在这种偏离并不引发进攻的假设下，这种偏离没有好处。如果偏离到策略 y 的做法招致布什进攻，这种偏离就更不利。再看类型 $\sim w$ 偏离均衡策略的潜在激励。在均衡中，他得到的收益为 2。如果他偏离均衡的行为没有招致布什进攻，他得到的效用为 3(也就是说，其境况因此改善)。因此，如果观察到对均衡策略的偏离，凭直觉判断，这种偏离最有可能来自 $\sim w$ 。Kreps 和 Cho 认为布什不会愚蠢地把 y 当作是萨达姆拥有杀伤性武器的证据。他们认为，对这种偏离路径行为的唯一合理解释是，类型为 $\sim w$ 的萨达姆可能会偏离到 $\sim w$ ，并发出了下面的信息：

亲爱的 W：

首先就过去与您父亲之间的瓜葛表示道歉。关于我们之间的分歧，我想做出解释：我知道你希望我选择 $\sim y$ 。但从这种选择中你无法做出任何推断——说白了，这是个混同均衡(你还记得落袋台球，对吧？你在耶鲁大学时玩过的)。但是，我不能这样做，因为我真的没有任何武器，在这一点上，我可以向全世界

保证。我现在的做法能给我带来更大好处。你应该知道,我采取的行动表明我真的没有任何武器,因为如果我藏有武器,并且我认为在我禁止武器核查时你不会进攻的话(这是我们正在使用的均衡策略),那么允许武器核查只会伤害我自己。如果我藏有武器,这种偏离对我没有任何好处。

真诚的

萨达姆·侯赛因

当然了,这种形式的沟通并没有出现在标准的信号揭示博弈中。它所表达的是,给定布什的策略,一种类型的萨达姆可能从偏离路径的选择中受益,另一种类型会蒙受损失。在这种背景下,偏离路径的推断应该集中到从中受益的类型上。这里要指出的是,如果混同均衡策略为 y ,就不存在这一问题。从选择 y 中可能受益的唯一类型是 w 。所以,在观察到 y 之后支持布什进攻的推断,是合理推断。

我们现在以更严格的方式阐述直觉标准,这要用到更多符号。设 Γ^S 为两阶段双人信号揭示博弈。参与者 1 有类型空间 Θ 和信号空间 M 。参与者 2 观察到参与者 1 的信号 m , 从 A 中选择一个行动。为简化分析,设所有这些集合都是有限集。对非有限集的更复杂的情况,我们只需调整下面的条件即可,但是会遇到一些技术性问题。参与者 1 知道自己的类型,参与者 2 只知道参与者 1 的类型来自 Θ 上的某个概率函数 $f(\cdot)$ 。参与者的收益分别由效用函数 $u_s(m, a, \theta)$ 和 $u_r(m, a, \theta)$ 决定。此时,混合策略组合由信号函数 $\sigma_s(\theta)$ 和行动函数 $\sigma_r(m)$ 构成;前者为每一个 θ 在 M 上选择一个概率分布,后者为每一个可能的信号选择行动集上的一个概率分布。设 $\sigma_s(m, \theta)$ 和 $\sigma_r(a, m)$ 分别表示类型为 θ 的信号发送者发出信号 m 的概率,及在观察到 m 后信号接收者 r 采取行动 a 的概率。序贯均衡还涉及推断 $\mu(\theta|m)$ 。给定信号揭示博弈和这一博弈的一个序贯均衡,设 $U_s^*(\theta)$ 表示类型为 θ 的参与者 1 从均衡策略组合中得到的期望效用。设 Δ 表示 Θ 上的概率分布集, $BR_r(m) = \bigcup_{p(\theta) \in \Delta} \{\arg\max_{a \in A} \sum u_r(m, a, \theta) p(\theta)\}$ 表示给定关于 θ 的某个推断时,最大化信号接收者的期望效用的行动所构成的集合。我们说行动 a 是可理性化的,如果它是 $BR_r(m)$ 中的元素。

定义 8.8 序贯均衡 $(\sigma_s(\cdot), \sigma_r(\cdot), \mu(\cdot|\cdot))$ 满足直觉标准,如果对满足 $\sum \sigma_r(m, \theta) f(\theta) = 0$ 的任意信号 m , 仅当 $U_s^*(\theta) < \max_{a \in BR_r(m)} u_s(m, a, \theta)$ 时,后验推断 $\mu(\theta|m) > 0$ 。

换句话说,直觉均衡(更准确地说,满足直觉标准的序贯均衡)要求,偏离均衡的推断赋予从所观察到的偏离中得不到任何好处的那些类型——当信号接收者通过使用其最优反应集中的策略对这种偏离做出反应时,得不到任何好处的那些类型——的概率为 0。

为了阐述相关概念,我们讨论前面的阻止进入博弈。这里不再限制信号空间为 $\{WC, \sim WC\}$, 而允许在位者自由选择资金筹措规模 $s_l \in \mathbb{R}_+$, 筹措成本为 cs_l 。根据在位者的类型, c 为 c_w 或 c_s 。设这一公职对在位者的价值为 1, 所以,如果在均衡中在位者筹措的资金规模为 s' 且以 π 的概率胜出,则在位者的收益是 $\pi - cs'$ 。在

这一博弈中,存在多个混同均衡、部分混同均衡和分离均衡。但是,只存在一个直觉均衡,它是个分离均衡。我们把对这一博弈的分析留做习题。

关于精炼的文献很多,对直觉标准的精炼出现在许多应用研究中。一般来说,类型空间的元素数量超过两个的模型,要求更强的精炼。这是因为,对接收方的各种最优反应,可能有多个类型能从偏离中受益。在这种情况下,更强的精炼,如**全域至圣均衡**(universal divinity)(Banks and Sobel, 1987),可能能够限制均衡集。^①由于全域至圣均衡被用于许多政治研究中,这里给出它的定义和一个例子。

定义 8.9 序贯均衡或 $PBE(\sigma_S(\cdot), \sigma_r(\cdot), \mu(\cdot|\cdot))$ 是全域至圣均衡,如果对满足 $\sum \sigma_r(m, \theta) f(\theta) = 0$ 的任意信号 m , 仅当存在行动 $a \in BR_r(m)$ 使得 $U_S^*(\theta) < u_s(m, a, \theta)$ 且对每个 $\theta' \neq \theta$, $U_S^*(\theta') \geq u_s(m, a, \theta')$ 时,后验推断 $\mu(\theta|m) > 0$ 。

全域至圣均衡比直觉标准更苛刻。它允许后验推断对更少数量的类型赋予正概率。在全域至圣均衡中,后验推断只对一种类型赋予权重,条件是存在可理性化的行动使得这种偏离对这种类型有利而对其他类型不利。不正式地说,全域至圣均衡要求偏离路径的推断只对“最可能”偏离的类型赋予权重。

为比较直觉标准和全域至圣均衡,我们把 Michael Spence 的信号揭示模型放到政治改革背景中(Spence, 1975)。假设某个发展中国家的类型为 $\theta \in \{1, 2, 3\}$, θ 表示这个国家成功偿付债务的能力(数字越大,能力越高)。设 π_1 和 π_2 分别表示这个国家为类型 1 和类型 2 的概率(它为类型 3 的概率是 $1 - \pi_1 - \pi_2$)。这个国家必须选择某一可被观察到的政治改革力度 $r \in \mathbb{R}_+$ 。改革成本决定于这个国家的类型。在观察到 r 后,IMF 决定为这个国家提供资金援助,数量为 $f \in \mathbb{R}_+$ 。信号接收者——即 IMF——的目标是把 f 与 θr 匹配。给定类型 θ 、改革 r 和援助 f , 这个国家的效用是 $f - \frac{r^2}{\theta}$ 。

首先设 $\pi_1 + \pi_2 = 1$ (就是说,只存在两种类型的发展中国家)。在这种情况下,存在混同 PBE、部分混同 PBE 和分离 PBE。但是,在直觉标准下只有一个均衡。我们来证明这一结论。图 8.11 中的信号—援助空间中的两条曲线分别为类型 1 和类型 2 的信号发送者的无差异曲线。所有的信号发送者都偏好更多的援助和更少的改革,所以,它们都希望无差异曲线向东北方向移动。

我们看一个混同均衡(或者说,部分混同均衡)。在这个均衡中,两种类型的发送者都以正概率选择相同的 r^p 。在观察到 r^p 后,信号接收者知道, $\theta < 2$ 的后验概率大于 0。因此,在任何 PBE 中,与 r^p 相对应的援助金额 $f(r^p)$ 肯定小于 $2r^p$ 。所以,存在信号 $r' > r^p$ 使得在 $f(r') = 2r'$ 时,类型为 2 的国家会偏离,类型为 1 的国家不会偏离。在这种情况下,根据直觉标准,在观察到 r' 时,所形成的推断赋予 $\theta = 2$ 的概率肯定是 1。给定这些推断,在观察到 r' 后,对信号接收者来说,金额为

① 建议读者阅读 Cho 和 Kreps(1987)与 Banks 和 Sobel(1987)。另外,Banks(1991)相当精彩地阐述了信号揭示博弈和精炼概念在政治科学中的应用。

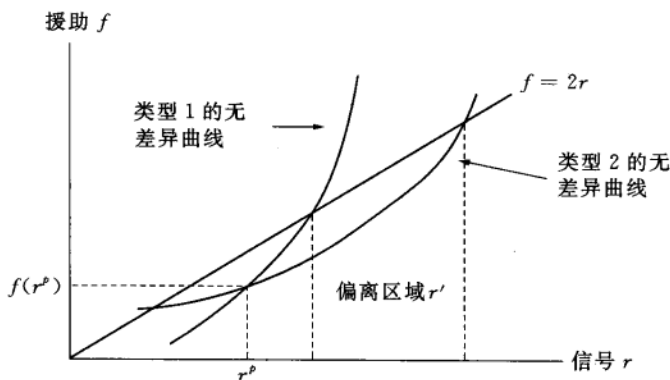


图 8.11 国际援助和改革博弈

$f(r') = 2r'$ 的援助是唯一的序贯均衡决策。于是就有：

$$U_r^*(2) = f(r^p) - \frac{(r^p)^2}{2}$$

$$u_r(r', 2r', 2) = 2r' - \frac{(r')^2}{2}$$

$$U_r^*(1) = f(r^p) - (r^p)^2$$

$$u_r(r', 2r', 1) = 2r' - (r')^2$$

因此, $U_r^*(1) > u_r(r', 2r', 1)$ 要求：

$$f(r^p) > (r')^2 - (r^p)^2$$

而 $U_r^*(2) < u_r(r', 2r', 2)$ 则要求：

$$f(r^p) < 2r' - \frac{(r')^2 - (r^p)^2}{2}$$

这两个不等式能同时得到满足。图 8.11 中给出了满足这两个不等式的 r' 的取值范围。简言之, 在混同均衡或者部分混同均衡中, 如果两种类型的国家都以正概率使用策略 r^p , 直觉标准要求观察到(可能偏离路径的)信号 r' 后, 所形成的推断赋予类型 $\theta = 2$ 的概率为 1。所以, 如果观察到 r' , 则资金援助为 $2r'$ 。对 r' 的值的选择不使得在信号 r' 和 r^p 之间, 类型为 2 的国家严格偏好 r' , 这意味着类型为 2 的国家不可能赋予信号 r^p 以正的概率。这与直觉均衡的假设, 即两种类型的国家都以正概率选择 r^p , 相矛盾。我们排除了分离均衡以外的所有均衡。但是, 我们还需要证明, 在这一博弈中, 直觉标准确定了唯一的分离均衡。作为习题, 我们请你证明这一结论。

假设 $\pi_1 + \pi_2 < 1$, 就是说, $\theta = 3$ 的概率为 $1 - \pi_1 - \pi_2 > 0$ 。这里要探讨的第一个问题是, 直觉标准是否依然能够剔除掉所有的混同均衡或部分混同均衡。答案是否定的。假设类型 1 和类型 2 都以正概率发出信号 r^p , 类型为 $\theta = 3$ 的国家则

使用纯策略 $r^3 > r^p$ 。我们能够证明在一些 PBE 中存在这样的策略。前面说过,直觉标准要求观察到改革 $r' > r^p$ 后,后验推断集中在类型高的国家上。但是,现在有了第三种类型,这种说法不再成立。假设在观察到改革 $r' > r^p$ 后,后验推断赋予 $\theta = 3$ 的权重满足直觉标准。在这样的推断下,对 f 的序贯理性选择可能导致类型 $\theta = 1$ 选择偏离而不会去获取均衡收益。具体地说,条件 $f(r') \leq 2r'$ 不再成立。现在,唯有 $f(r') \leq 3r'$ 成立。在 $f(r') \geq 3r'$ 时,类型 $\theta = 1$ 预期自己的偏离会导致 $f(r')$, 可能会选择偏离。类型 $\theta = 2$ 为表明自己不是类型 1, 所发出的信号至少必须和图 8.12 中的 $r^{\min}$ 一样高。但是,对高于 $r^{\min}$ 的所有的 r , 可能存在最优反应 $f(r)$ 使得类型 2 得到的效用低于均衡收益。从图上看,在类型 2 的无差异曲线和直线 $f = 2r$ 之间存在一块区域。由于直觉标准只要求对大于 $r^{\min}$ 的 r , 类型 2 能得到的 IMF 的反应为 $f > 2r$, 所以,这种偏离是否有利是不确定的。

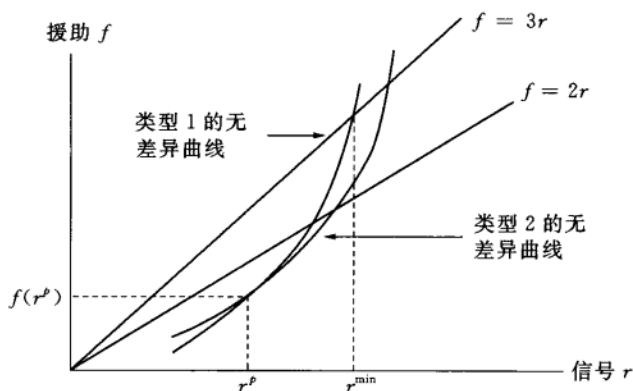


图 8.12 国际援助和改革博弈

但是,这个部分混同均衡不可能是全域至圣均衡。假设类型 1 和类型 2 都以正概率发出信号 r^p , 类型 $\theta = 3$ 使用纯策略 $r^3 > r^p$ 。在全域至圣均衡中,略高于 r^p 的信号 r' 肯定导致集中在 $\theta = 2$ 上的后验推断; 当 r 位于 r^p 右边时,类型 1 的无差异曲线在类型 2 的无差异曲线上方。对位于这两条无差异曲线之间的 $(r', f(r'))$, $U_2^*(r') < u_2(r', f(r'), 2)$ 且 $U_1^*(r') > u_1(r', f(r'), 1)$ 。而且,由于类型 3 得到 $f = 3r$, 所以和偏离相比,它在均衡中得到的效用更高。就是说,全域至圣均衡要求信号 r' 导致推断集中在 $\theta = 2$ 上,序贯理性要求 $f(r') = 2r'$ 。类型 2 从偏离中得到好处。在这个有三种类型的博弈中,只存在一个全域至圣均衡。作为习题,请你证明这一结论。

8.7 习题

习题 8.1 在图 8.6 的博弈中,设路径 (B, NR) 的收益是 $(5, 0)$ 而不是 $(4, 0)$ 。找出这一博弈的所有的(混合策略和纯策略)PBE。

习题 8.2 在图 8.6 的博弈中, 设路径 (N, D) 的收益是 $(w, 5)$ 而不是 $(0, 5)$ 。这里的 w 是外生参数, 参与者知道它在 -2 到 5 之间取值。在 w 取什么值时, 存在 PBE, 且在这一均衡中, (N, D) 以正概率发生?

习题 8.3 证明在图 8.7 的博弈中, 不存在 $m(\theta) = \begin{cases} b, & \text{如果 } \theta = a \\ a, & \text{如果 } \theta = b \end{cases}$ 的 PBE。

习题 8.4 找出图 8.13 的博弈中的所有 PBE。

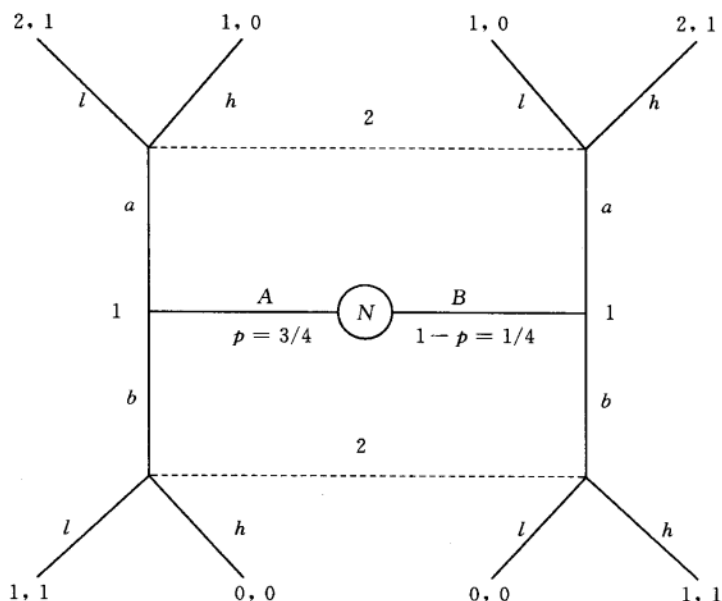


图 8.13 信号揭示博弈

习题 8.5 政治压迫模型: 假设在两个阶段中的每个阶段里, 社会 (S) 必须决定是否反对政府的政策。政府可以选择接受或者压制。政府如果接受其反对意见, 社会得到 1 个单位的收益; 政府如果压制其反对意见, 社会得到的收益为 -1 ; 如果没有提出反对, 社会得到 0 个单位的收益。

假设有两种类型的政府: 温和型和强硬型。如果没有来自社会的反对, 温和型政府 (M) 得到的收益为 0; 如果存在社会反对意见而政府接受社会反对意见, 温和型政府得到的收益为 -2 ; 如果它压制社会的反对意见, 得到的收益为 -3 。如果没有来自社会的反对, 强硬型政府 (H) 得到的收益为 0; 如果它压制社会的反对意见, 得到的收益为 -2 ; 如果它接受社会的反对意见, 得到的收益为 -3 。设 p_0 是政府为类型 M 的先验概率。

(1) 在第二个阶段里, 设在 $p_1 \geq p^*$ 时 (其中, p_1 是 S 对政府为类型 M 的后验推断), S 选择反对政府政策。请问, 临界值 p^* 是多少?

(2) 是否存在包含以下策略的分离均衡?

M : {接受, 接受}

$H: \{\text{压制, 压制}\}$

$S: \{\text{在第一阶段反对, 如果在第一个阶段被压制, 则在第二阶段不再反对}\}$

如果存在这样的分离均衡, 则 p_0 的值必须为多少?

(3) 是否存在使用以下策略的第一阶段混同均衡?

$M: \{\text{压制, 接受}\}$

$H: \{\text{压制, 压制}\}$

$S: \{\text{反对, 如果在第一个阶段中被压制, 则在第二个阶段不再反对}\}$

如果存在这样的混同均衡, 则 p_0 的值必须为多少? 它与直觉标准相一致吗?

(4) 是否存在使用以下策略的第一个阶段混同均衡?

$M: \{\text{压制, 接受}\}$

$H: \{\text{压制, 压制}\}$

$S: \{\text{不反对, 如果在第一个阶段中被压制, 则在第二阶段不再反对}\}$

如果存在这样的混同均衡, 则 p_0 的值必须为多少? 它与直觉标准相一致吗?

(5) 计算半混同均衡, 在这一均衡中, M 在第一个阶段中以概率 q 选择压制。在 p_0 为多少时, S 反对政府政策呢?

习题 8.6 求解竞选筹资博弈中其他的部分混同均衡。

习题 8.7 在 Gilligan-Krehbiel 模型中, 证明在开放规则下不存在部分混同均衡。

习题 8.8 在 Gilligan-Krehbiel 模型中, 在开放规则下, 假设委员会中只有两个成员, 其理想点为 $c > \theta$ 。他们观察到真实状态, 同时向议会发出信号。在这一博弈中, 存在分离均衡吗? 如果委员会中有三个成员, 存在分离均衡吗?

习题 8.9 对任意的有限扩展式博弈, 如果 $(\sigma(\cdot), b(\cdot))$ 构成序贯均衡, 证明 $\sigma(\cdot)$ 是子博弈精炼。

习题 8.10 在只有一个信号发送者和一个信号接收者的信号揭示博弈中, 如果发送者的类型空间中有两个元素, 则 PBE 集和 SE 集相同。

习题 8.11 证明命题 8.9。

习题 8.12 证明命题 8.10。

习题 8.13 在萨达姆—布什博弈中, 证明使用策略 y 的混同均衡没有违背直觉标准。

习题 8.14 通过调整图 8.10 中的概率 w , 能够使所观察到的行动路径 $(\sim y, a)$ 构成 PBE 吗?

习题 8.15 在竞选筹资博弈中, 设信号空间为正实数空间。分析能实现的 PBE 的 s_i 的大小和进入行为上的概率。

习题 8.16 在竞选筹资博弈中, 设信号空间为正实数空间。求这一博弈的唯一的直觉均衡。

习题 8.17 在贷款担保博弈中, 设只有两种类型的借款国。证明在唯一的直觉均衡中:

$$r(1) = \operatorname{argmax} \{r - r^2\} = \frac{1}{2} \text{ 且 } r(2) = \operatorname{argmax} \left\{r - \frac{r^2}{2}\right\} = 1$$

习题 8.18 在贷款担保博弈中, 设有三种类型的借款国。证明在唯一的全域至圣均衡中:

$$r(1) = \operatorname{argmax} \{r - r^2\} = \frac{1}{2},$$

$$r(2) = \operatorname{argmax} \left\{r - \frac{r^2}{2}\right\} = 1,$$

$$r(3) = \operatorname{argmax} \left\{r - \frac{r^2}{3}\right\} = \frac{3}{2}.$$

新
解
題
PDG

► 9

重复博弈

在政治博弈论的许多模型中,参与者重复进行同一个博弈。其中一些模型研究社会习俗,如惯例、规范、合作和信任等,在参与者有偏离这些习俗的短期激励时,是如何得以持续存在的。“重复博弈”的另一重要应用是研究在不求助中央权力时,如何解决社会两难问题,如囚徒两难等(Taylor, 1976)。

在这种博弈中,最有意义的概念性问题是,相对于一次性博弈,这种重复性在多大程度上使得更多的行为成为 Nash 均衡行为? 一般情况下,在重复博弈中,相对于它所对应的静态博弈, Nash 均衡集大很多。未来的行为依赖于当前的行为。对行为的这种预期能为一次性博弈中的非最优行为创造激励。但是,重复博弈又有另一方面的问题;它所包含的均衡如此之多,很难根据均衡准确预测博弈结果。

我们探讨关于未来的预期如何影响行为。我们探讨表 9.1 中的标准式博弈。^①

这个博弈如果只进行一次,则有两个 Nash 均衡: (M, M) 和 (B, R) (表中加星号的单元格)。策略组合 (T, L) 提供最高总收益,但不是 Nash 均衡;参与者 1 会单方面偏离到 B , 参与者 2 会单方面偏离到 R 。在这一博弈进行两次并且两个参与者都只关心自己在两个时期中的总收益时,结果又如何呢? 假设参与者 1 选择的策略是“在时期 1 使用 T ; 如果参与者 2 在时期 1 使用 L , 则在时期 2 使用 M ; 否则,就在时期 2 使用 B ”。假设参与者 2 选择的策略是“在时期 1 使用 L ; 如果参与者 1 在时期 1 使用 T , 则在时期 2 使用 M ; 否则,就在时期 2 使用 R ”。如果上述阶段博弈只进行两次,这对策略就构成 Nash 均衡。在这个均衡中,两个参与者在第一个时期里实现“好的”结果 (T, L) 。如果有人偏离了这一结果,他的收益就为 $9+3=12$, 小于均衡效用 $8+5=13$ 。事实上,这个策略不仅构成 Nash 均衡,而且还是子博弈精炼均衡。 (M, M) 和 (B, R) 都是一次性博弈的 Nash 均衡,所以在真子博弈中使用这样的策略组合自然与子

表 9.1

1\2	<i>L</i>	<i>M</i>	<i>R</i>
<i>T</i>	8, 8	0, 0	1, 9
<i>M</i>	0, 0	5, 5*	0, 0
<i>B</i>	9, 1	0, 0	3, 3*

^① 到了这里,我们觉得已不再需要用例子阐述博弈理论,想来读者也能接受这一点。

博弈精炼相一致。但是,在两个时期使用(T, L)的策略组合却不会出现在 Nash 均衡中。

9.1 重复的囚徒两难博弈

研究得最多的博弈之一是重复的囚徒两难博弈。我们看它在贸易政策上的应用。如果所有国家都实行自由贸易,世界经济会运转地更好,但是,各个国家都想

表 9.2 自由贸易博弈

美国\欧盟	自由贸易 F	保护 P
自由贸易 F	10, 10	1, 12
保护 P	12, 1	4, 4

保护本国经济。在存在着这种冲突时,自由贸易机制如何实现? 一种答案是,自由贸易可能是重复博弈中的均衡——在这一博弈中,只要一个大国偏离贸易协定,就会引发贸易战。为阐述这一点,我们看表

9.2 中的博弈,它描述了美国(US)和欧盟(EU)在贸易政策上的两难。

显然,这一博弈如果只进行一次,唯一的 Nash 均衡是策略组合(保护,保护)。如果它进行两次,则每个参与者的策略集是

$\{FT_1FT_2FT_2, FT_1FT_2P_2, FT_1P_2FT_2, FT_1P_2P_2, P_1FT_2FT_2, P_1FT_2P_2, P_1P_2FT_2, P_1P_2P_2\}$

其中, $FT_1FT_2P_2$ 说的是“在时期 1 使用 FT ,如果对方在时期 1 使用 FT ,则在时期 2 使用 FT ;否则,就使用 P ”。表 9.3 给出了这个两时期博弈的标准形式(这里不考虑未来收益的贴现)。

表 9.3 两阶段自由贸易博弈

美国\欧盟	$FT_1FT_2FT_2$	$FT_1FT_2P_2$	$FT_1P_2FT_2$	$FT_1P_2P_2$	$P_1FT_2FT_2$	$P_1FT_2P_2$	$P_1P_2FT_2$	$P_1P_2P_2$
$FT_1FT_2FT_2$	20, 20	20, 20	11, 22	11, 22	11, 22	11, 22	2, 24	2, 24
$FT_1FT_2P_2$	20, 20	20, 20	11, 22	11, 22	13, 13	13, 13	5, 16	5, 16
$FT_1P_2FT_2$	22, 11	22, 11	14, 14	14, 14	11, 22	11, 22	2, 24	2, 24
$FT_1P_2P_2$	22, 11	22, 11	14, 14	14, 14	13, 13	13, 13	5, 16	5, 16
$P_1FT_2FT_2$	22, 11	13, 13	22, 11	13, 13	14, 14	5, 16	14, 14	5, 16
$P_1FT_2P_2$	22, 11	13, 13	22, 11	13, 13	16, 5	8, 8	16, 5	8, 8
$P_1P_2FT_2$	24, 2	16, 5	24, 2	16, 5	14, 14	5, 16	14, 14	5, 16
$P_1P_2P_2$	24, 2	16, 5	24, 2	16, 5	16, 5	8, 8	16, 5	8, 8

和表 9.1 中的例子不同,把这一博弈重复一次并不影响参与者的行为,因为($P_1P_2P_2, P_1P_2P_2$)是唯一的 Nash 均衡。这一结论适用于任何有限个时期博弈。在最后一个时期,每个国家都推行保护政策。在倒数第二个时期,每个参与者都认

识到了这一点,所以,倒数第二个时期中的行为不可能影响最后一个时期的行为。因此,在倒数第二个时期里,每个国家都有动机推行保护政策。这一过程一直持续下去,直至每个国家在每个时期里都采取保护政策。

在表 9.1 中,我们可以诱生出时期 1 中的合作,因为时期 1 的行为有助于协调第二个时期中的多个均衡。好的均衡被用作奖励,差的均衡被用作惩罚。在囚徒两难博弈中,由于只有一个 Nash 均衡,所以无法通过协调到好的均衡上或者通过协调到差的均衡上,鼓励合作。

但是,如果这一博弈持续无限个时期,这一问题就不再是问题。假设“好的均衡”是在每个时期中自由贸易,“差的均衡”是在每个时期中贸易保护。我们将看到,因为不存在最后一个时期,好的均衡未必像在有限情况中那样不存在。就是说,在每个时期中,把实现好的均衡作为奖励和把实现差的均衡作为惩罚,合作得以持续。

9.2 触发均衡

为说明无限次重复如何消除了这“最后一个时期”的问题,我们分析这样一个策略:“在每个时期推行自由贸易直至对方实行贸易保护,此后始终实行贸易保护。”这就是触发策略,因为只要偏离合作,就会导致未来所有时期里的非合作均衡。如果每个国家都采取这一策略,在每个时期里双方都得到 10 个单位的效用。两个国家如果以相同的贴现因子 δ 贴现未来,这一策略实现的长期效用是 $\frac{10}{1-\delta}$ 。^①要证明这一策略是 Nash 均衡,必须证明两个参与者都不会选择偏离。在这一重复博弈中,由于所有的阶段博弈相同,所以,参与者或者在每个时期中都选择偏离,或者永远不选择偏离。在发生偏离的阶段里,偏离者得到 12 个单位的效用。由于在偏离之后,对方永远推行保护政策,所以,偏离者的最优反应是一直偏离下去。因此,在偏离之后的每个时期中,偏离者得到 4 个单位的效用。他从偏离行为中得到的总效用是 $12 + \delta \frac{4}{1-\delta}$ 。因此,在这一重复进行的囚徒

两难博弈中,当且仅当 $\frac{10}{1-\delta} \geq 12 + \delta \frac{4}{1-\delta}$, 触发策略构成 Nash 均衡。当且仅当

$\delta \geq \frac{1}{4}$ 时,这一充分必要条件成立。就是说,只要参与者足够地有耐性(δ 很大),触发策略就是 Nash 均衡。它也是子博弈精炼均衡。这个博弈的真子博弈也是无限次重复的囚徒两难博弈;所以,在这些子博弈中,只要此前有一方实行保护政策,永久的贸易保护就是每个子博弈的 Nash 均衡。而从双方皆推行自由贸易开始的子博弈,和原始博弈完全相同。由于这一策略组合是 Nash 均衡,所以它也是这一

① 关于时间贴现的讨论和贴现效用的无限和的计算,见第 3 章。

类型的子博弈中的 Nash 均衡。

表 9.4 一般的囚徒两难博弈

1\2	合 作	不合作
合 作	a, a	d, c
不合作	c, d	b, b

再看表 9.4 中的一般形式的囚徒两难博弈。其中, $c > a > b > d$ 。利用和上面的思路, 当且仅当:

$$\frac{a}{1-\delta} \geq c + \delta \frac{b}{1-\delta}$$

触发策略构成 SPNE。移项后, 得到:

$$\delta \geq \frac{c-a}{c-b}$$

就是说, 当:

- (1) c 相对于 a 和 b 很大时,
- (2) a 和 b 大致相等时,

合作更能维持下去(要求有更大的贴现因子)。

9.3 以牙还牙策略

触发策略并不是这一无限次重复进行的囚徒两难博弈中实现合作结果的唯一均衡。许多人认为, 触发均衡并不现实, 也不合理, 因为在一次背叛之后合作就永久地消失了。而且, 它不允许参与者犯错误。并且, 在合作破裂后, 参与者无法通过谈判回到合作阶段, 而他们显然都有动机寻求再次合作。但是, 如果参与者能够进行这种谈判, 合作的激励就不复存在; 非合作路径不再是可信赖的威慑力量。

另一种 Nash 均衡建立在“以牙还牙”策略上, 其形式为“在第一个时期合作, 在此后每个时期里使用前一个时期里对方所使用的行动”。为证明以牙还牙策略构成 Nash 均衡, 我们分析参与者 i 的单方面背叛: 在时期 t , 他选择不合作。在时期 t , 他的收益为 c 。在时期 $t+1$, 根据他所采取的行动, 他得到 d 或 b 。在时期 $t+1$ 回到以牙还牙策略以后, 他所实现的收益流为 $d, c, d, c, d, c, \dots$ 。所以, 在某个时期背叛和再回到以牙还牙策略上, 所实现的收益为:

$$c + \delta d + \delta^2 c + \delta^3 d + \delta^4 c + \dots = c + \delta d + \delta^2(c + \delta d) + \delta^4(c + \delta d) + \dots = \frac{c + \delta d}{1 - \delta^2}$$

如果自时期 t 开始的所有时期里, 参与者 i 始终不合作, 他的收益为 $c + \frac{\delta b}{1 - \delta}$ 。作为习题, 请你证明: 对以牙还牙策略的最优偏离, 所实现的收益是这两个收益流中的一个。前面说过, 均衡收益流为 $\frac{a}{1 - \delta}$ 。所以, 当且仅当:

$$\frac{a}{1 - \delta} \geq \max \left\{ \frac{c + \delta d}{1 - \delta^2}, c + \frac{\delta b}{1 - \delta} \right\}$$

以牙还牙策略构成 Nash 均衡。通过调整,我们就能将这个不等式表示为贴现因子所需满足的条件。不等式 $\frac{a}{1-\delta} \geq \frac{c+\delta d}{1-\delta^2}$ 可以简化为:

$$\delta \geq \frac{c-a}{a-d}$$

不等式 $\frac{a}{1-\delta} \geq c + \frac{\delta b}{1-\delta}$ 可以简化为:

$$\delta \geq \frac{c-a}{c-b}$$

就是说,当且仅当:

$$\delta \geq \max\left\{\frac{c-a}{a-d}, \frac{c-a}{c-b}\right\}$$

以牙还牙策略构成 Nash 均衡。

以牙还牙策略是子博弈精炼吗? 在参与者使用这一策略时,我们需要分析两类子博弈:

(1) 在前一个时期的合作之后,两个参与者在接下来的一个时期里合作。这是合作阶段。

(2) 在一个参与者选择合作而另一个参与者选择背叛之后,背叛者选择合作而对方在接下来的一个时期里选择不合作来惩罚背叛者。这是惩罚阶段。

事实上,在这个博弈中还存在一种子博弈,其中每个人都选择背叛。但是在双方都使用以牙还牙策略时,对均衡路径的单方面偏离不会到达这种子博弈。我们借助单偏离原则来理解合作阶段中的偏离动机。此前在探讨子博弈精炼时,我们证明对有限博弈,单偏离原则成立。对无限重复博弈,在收益以贴现率 $\delta < 1$ 贴现时,这一原则也成立。我们把这一结论的证明留做习题。^①对以牙还牙策略的单方面偏离,所实现的收益流是 $c + \delta d + \delta^2 c + \delta^3 d + \delta^4 c + \dots$, 前面的分析指出,如果 $\delta \geq \frac{c-a}{a-d}$, 这种偏离就得不偿失。

我们看惩罚阶段。假设在以牙还牙策略下,参与者 i 应该选择不合作,参与者 j 应该选择合作。参与者 i 在时期 t 的单方面偏离,导致双方在时期 t 以后的所有时期里都选择合作。于是, i 从这次偏离中得到的收益是 $a + \delta a + \delta^2 a + \dots = \frac{a}{1-\delta}$ 。而在惩罚阶段中使用均衡策略所实现的收益流是 $c + \delta d + \delta^2 c + \delta^3 d + \dots = \frac{c+\delta d}{1-\delta^2}$ 。因此,只要 $\frac{c+\delta d}{1-\delta^2} \geq \frac{a}{1-\delta}$, 在惩罚阶段里参与者 i 会使用以牙还牙策略。

① 把单偏离原则推广到无限时域博弈中的困难是,不构成子博弈精炼的策略可能会输给持续无限个时期的偏离。但是,只要能对未来贴现,在足够远的时期中发生的偏离就不可能对收益有很大影响,因此,如果一个策略不是子博弈精炼,肯定就存在更好的策略,它包含只持续有限个时期的偏离。从这一思路出发,有限博弈中的得到结论适用于无限博弈。

这一不等式与以牙还牙策略构成 Nash 均衡的条件正好相反。惩罚策略要奏效,相对于由 c 构成的持久收益流, i 必须偏好在 c 和 d 之间来回震荡的收益流;在以牙还牙策略和开始之时便偏离的策略之间,要使 i 偏好前者,则在由 a 构成的持久收益流和振荡收益流之间,他就必须偏好前者。就是说,只在 $\delta = \frac{c-a}{a-d}$ 时——一个刀锋状条件——以牙还牙策略才构成子博弈精炼 Nash 均衡。

另一种以牙还牙策略则避免了惩罚阶段中的振荡问题。Milgrom, North 和 Weingast(1990)谈到了一种策略:“开始时选择合作。除非你在时期 $t-2$ 选择不合作,否则,当对手在时期 $t-1$ 选择合作时,你在时期 t 选择合作。如果你在时期 $t-1$ 选择不合作,则在时期 t 和时期 $t+1$ 都选择合作。”这一策略虽然有些复杂,却有很直观的含义。它指出,在任何偏离之后,偏离方必须在下个时期中,当对手选择背叛时选择合作。使用这样的策略,就能用低收益 d 惩罚偏离者。在这个单时期惩罚之后,双方回到合作状态中。作为习题,请你证明:只要 δ 足够大,这种以牙还牙策略是子博弈精炼。

9.4 介于触发策略和以牙还牙策略之间的惩罚策略

触发策略和以牙还牙策略只是实现合作结果的两种可能的策略。我们可以把这两种策略推广开来,允许惩罚阶段的时间长度在两者之间。我们看下面的策略:

(1) 合作,直至对手背叛。如果对手背叛,则在接下来的 k 个时期中选择不合作。在 k 个时期过后,开始合作。如果你选择背叛,则在接下来的 k 个时期依然选择不合作,在 k 个时期过后再回到合作状态。一旦回到合作状态,就一直合作下去,直至再次发生背叛。在发生背叛后,再次使用上述惩罚策略。

(2) 合作,直至对手背叛。如果对手背叛,则在接下来的 k 个时期中不合作。如果在这 k 个时期的每个时期中,对手选择合作,就结束惩罚阶段,开始合作。如果在惩罚阶段中的任何一个时期里,对手都没有选择合作,则开始新一轮惩罚阶段,就是说,在接下来的 k 个阶段中,依然选择不合作。如果由于你自己没有选择合作而进入了惩罚阶段,则在惩罚阶段中选择合作。

策略 1 类似触发策略,因为,惩罚由一对策略(不合作,不合作)构成。只是在这里,惩罚阶段由有限个时期构成。策略 2 类似以牙还牙策略的子博弈精炼形式,因为背叛者被惩罚是出于合作的需要。

首先探讨策略 1。在惩罚阶段中,双方没有偏离这一策略的动机:就当期收益而言,选择不合作是对对手的不合作行为的最优反应,偏离却只会延长惩罚时间。假设在双方合作时发生了背叛。合作阶段中的一次背叛所实现的收益流包括:从背叛行为中获得的单时期收益、在接下来的 k 个时期中每个时期得到收益 b 、自未来第 $k+1$ 个时期开始的每个时期为 a 的无限收益流。利用子博弈精炼的

单偏离原则,我们只需分析一次偏离所导致的后果。利用无限和的计算法则,惩罚阶段中的背叛行为所实现的效用为:

$$c + \frac{\delta - \delta^{k+1}}{1 - \delta} b + \frac{\delta^{k+1}}{1 - \delta} a$$

因此,要维持合作,要求:

$$a(1 - \delta^{k+1}) \geq (1 - \delta)c + (\delta - \delta^{k+1})b$$

对 δ 的临界值,我们无法得到闭式解,但是,我们可以把这一表达式写为:

$$\delta > \frac{c - a}{c - b} + \delta^{k+1} \frac{a - b}{c - b}$$

不等式右边的第一项是触发策略的临界值;对任何有限的 k ,第二项是正的。不难看出,在惩罚期由有限个时期构成时,实现合作的难度更大。但是,在参与者可能犯错误的模型中,这一均衡可能优于触发策略。^①

现在探讨策略 2。假设在合作阶段中发生背叛。背叛行为所实现的收益包括:单时期收益 c ,连续 k 个时期的惩罚收益 d ,在惩罚期后恢复合作所实现的收益 a 。将这些收益加到一起,得到 $c + \frac{(\delta - \delta^{k+1})d}{1 - \delta} + \frac{\delta^{k+1}a}{1 - \delta}$ 。由于 $a > d$,所以这一表达式在 k 上递减。因此,增加惩罚期的长度,削弱在合作阶段中发生背叛行为的激励。

但是,增加 k 可能不会使这种均衡更易实现,因为它削弱了在惩罚阶段中使用这一策略的激励。我们分析在惩罚阶段中发生的背叛行为所实现的收益,它包括:单时期收益 b ,连续 k 个时期获得收益 d ,然后开始获得收益 a 。总收益为 $b + \frac{(\delta - \delta^{k+1})d}{1 - \delta} + \frac{\delta^{k+1}a}{1 - \delta}$ 。在惩罚阶段中坚持上述均衡策略所实现的收益,则取决于双方现在位于惩罚阶段中的哪个时期上。要证实双方在每个时期里坚持均衡策略,只需证实在均衡收益最低的时期中,双方坚持均衡策略。这个时期便是惩罚阶段中的第一个时期。在这个时期中,坚持惩罚策略所实现的效用为 $\frac{(1 - \delta^k)d}{1 - \delta} + \frac{\delta^k a}{1 - \delta}$ 。因此,要坚持这种惩罚策略,要求:

$$\delta > \left(\frac{b - d}{a - d} \right)^{\frac{1}{k}}$$

这一临界值显然在 k 上递减。

① 和使用触发策略的 SPNE 一样,这种惩罚策略无法抵挡通过谈判达成和解的诱惑。在双方进入到惩罚阶段后,如果他们能通过谈判达成妥协,在少于 k 个时期后就恢复合作,双方都能从中得到好处。

9.5 无名氏定理

上述例子的共同点是,只要参与者有足够的耐性,在静态博弈中不构成 Nash 均衡的结果,可能是无限重复博弈的 Nash 均衡或者子博弈精炼均衡。事实上,在无限重复博弈中,只要参与者有足够的耐性,满足个人理性的任何收益向量都可能构成 SPNE。这一结论很久以来就被博弈学家们所了解,并因此得到了无名氏定理的称号。在本节中,我们证明一种形式的无名氏定理。

重复博弈的构成要素包括标准式阶段博弈 $\Gamma = \langle N, S, u \rangle$ 和参与者的贴现因子向量 $\delta = (\delta_1, \dots, \delta_n)$ 。在每个时期 $t \in \{1, 2, 3, \dots\}$ 中,参与者进行标准式博弈 Γ 。在参与者 i 在时期 t 选择自己的策略 $s_i^t \in S_i$ 之前,他观察到时期 $t-1$ 中所使用的策略组合 s^{t-1} 。我们假定存在完美记忆(perfect recall),于是, s_i^t 就决定于历史 $h^{t-1} = (s^1, \dots, s^{t-1}) \in S^{t-1} := \prod_{j=1}^{t-1} S_j$ 。空历史为 $h^0 = \emptyset$ 。参与者 i 的一个纯策略于是是一个映射列 $\{s_i^t(h^{t-1}): S^{t-1} \rightarrow S_i\}_{t=1}^\infty$; 一个混合策略就是一个映射列 $\{\sigma_i^t(h^{t-1}): S^{t-1} \rightarrow \Delta(S_i)\}_{t=1}^\infty$ 。给定阶段博弈上的概率分布组合构成的列 $\{\sigma_i\}_{i=1}^\infty$, 参与者 i 的期望效用为 $EU_i(\{\sigma_i\}_{i=1}^\infty) = (1 - \delta_i) \sum_{t=1}^\infty \delta_i^{t-1} E_{\sigma^t} u_i(s^t)$, 其中 $E_{\sigma^t} u_i(s^t)$ 是 $u_i(s^t)$ 在 σ^t 上的期望。引入乘数 $(1 - \delta_i)$ 。对常数列 σ^t , $EU_i(\{\sigma_i\}_{i=1}^\infty) = Eu_i(\sigma^t)$ 。由阶段博弈构成的重复博弈被表示为 $\Gamma^\infty = \langle N, S, u, \delta \rangle$ 。由于重复博弈也是扩展式博弈,因此,在重复博弈中, Nash 均衡和子博弈精炼 Nash 均衡等概念都有定义。

我们探讨由有限标准式阶段博弈生成的重复博弈。根据 Nash 均衡,这种阶段博弈至少有一个混合策略 Nash 均衡。因此不难想象,这种重复博弈有混合策略子博弈精炼 Nash 均衡;不管所发生的历史是什么,在每个时期中都使用阶段博弈的均衡策略的策略,构成 Γ^∞ 的 SPNE。

定理 9.1 如果 σ^* 是阶段博弈的子博弈精炼 Nash 均衡,则对每个 i , 对每个 t , 对每个 h^{t-1} , 重复博弈的策略组合 $\sigma_i^t(h^{t-1}) = \sigma_i^*$ 是这一重复博弈的子博弈精炼 Nash 均衡。

在重复博弈中,子博弈精炼 Nash 均衡集通常很大。无名氏定理能够确定在这种均衡中可能实现的均衡收益集。我们证明一个十分有用的简单的无名氏定理。我们先给出几个定义。

定义 9.1 收益向量 $v = (v_1, \dots, v_i, \dots, v_n) \in \mathbb{R}^n$ 满足个人理性,如果对每个 $i \in N$,

$$v_i \geq \min_{s_{-i} \in S_{-i}} \{ \max_{s_i \in S_i} u_i(s_i, s_{-i}) \}$$

其中, $\min_{s_{-i} \in S_{-i}} \{ \max_{s_i \in S_i} u_i(s_i, s_{-i}) \}$ 是对手使用 s_{-i} 时,参与者 i 使用自己的最优反应所能获得的最小的阶段博弈收益。参与者 $-i$ 通过选择 s_{-i} , 最小化参与者 i 通过使用对 s_{-i} 的最优反应所获得的效用,就能够得到这个值。

定义 9.2 收益向量 $v \in \mathbb{R}^n$ 是可行的, 如果存在某个纯策略组合 s , 使得对每个 $i \in N$, $\frac{u_i(s)}{1 - \delta_i} = v_i$ 。

定理 9.2 对每个可行的和满足个人理性的收益向量 $v \in \mathbb{R}^n$, 存在 n 维贴现因子向量 δ' , 在 δ 的每个坐标至少和 δ' 中的相应坐标一样大时, 收益向量 v 就给出了这个重复博弈的某个 Nash 均衡中的收益。

证明: 假设 v 是可行的且满足个人理性。 $\{s^t\}$ 表示一个策略组合——只要没有任何人此前偏离过这一策略或者没有两个以上的参与者曾经偏离过这一策略, 这个策略组合就能实现收益向量 v 。如果只有一个参与者, 设参与者 i , 曾经偏离过这一策略, 这一策略就要求:

$$\{s_i^t = \operatorname{argmin}_{s_{-i} \in S_{-i}} \{\max_{s_i \in S_i} u_i(s_i, s_{-i})\}\}$$

这一策略在随后的所有时期中惩罚偏离者 i 。在任何时期 t , 参与者 i 从使用 $\{s^t\}$ 中得到的收益为 v_i 。他从偏离行为中得到的收益的下界为:

$$(1 - \delta_i') v_i + \delta_i' (1 - \delta_i) \max_{s \in S} u_i(s) + \delta_i'^{t+1} \min_{s_{-i} \in S_{-i}} \{\max_{s_i \in S_i} u_i(s_i, s_{-i})\}$$

如果:

$$\delta_i \geq \frac{\max_{s \in S} u_i(s) - v_i}{\max_{s \in S} u_i(s) - \min_{s_{-i} \in S_{-i}} \{\max_{s_i \in S_i} u_i(s_i, s_{-i})\}}$$

这个下界值就小于 v_i 。由于 $\max_{s \in S} u_i(s) \geq v_i \geq \min_{s_{-i} \in S_{-i}} \{\max_{s_i \in S_i} u_i(s_i, s_{-i})\}$, 所以上面的不等号右边部分严格小于 1。只要这一条件对每个 $i \in N$ 都成立, 我们所猜想的策略组合就是这个重复博弈的 Nash 均衡。□

这个证明中所谈到的均衡未必是 SPNE, 因为其中的惩罚措施的成本可能很高, 因此无法实施。我们可以用阶段博弈的 Nash 均衡策略, 作为惩罚措施, 找到重复博弈的 SPNE 所支持的收益向量集。作为习题, 请你证明下面的结论。

定理 9.3 如果 $v \in \mathbb{R}^n$ 是可行的收益向量, 并且对每个 $i \in N$ 存在阶段博弈的某个混合策略 Nash 均衡, 它产生的期望收益向量为 v_i' , 其中 $v_i' < v_i$, 则存在 n 维贴现因子 δ' 使得在重复博弈中存在子博弈精炼 Nash 均衡, 如果 δ 的每个坐标至少与 δ' 中相对应坐标一样大, 这个均衡就产生收益向量 v 。

无名氏定理在政治科学中的大部分应用, 都没有超出上述两个命题。但是, 学者们把这两个结论在多个方向上进行了发展, 感兴趣的读者可以参考 Abreu (1988), Fudenberg 和 Maskin (1986)。

9.6 应用: 团体间合作

Fearson 和 Laitin (1996) 用无限重复博弈分析团体之间的合作如何实现。有两个团体 A 和 B, 它们各有 n 个成员 (n 为偶数)。在每个时期 t , 参与者们被随机

配对,进行表 9.5 中的囚徒两难博弈。在表 9.5 中, $\alpha > 1$, $\beta > 0$, $\frac{\alpha - \beta}{2} < 1$ 。假设

表 9.5 团体之间的合作博弈

$1 \backslash 2$	合作	背叛
合作	1, 1	$-\beta, \alpha$
背叛	$\alpha, -\beta$	0, 0

每个团体成员有相同的贴现因子 $\delta \in (0, 1)$ 。在每个时期里,每个团体中有 m 个成员被选中与另一团体中的成员配对,剩下的 $n - m$ 个成员与自己团体中的成员配对。在这一随机配对过程中,每个参与者有 $p = m/n$ 的概率与“外族”成员配对。

为探讨团体间关系和团体内关系的动态性, Fearon 和 Laitin 假定,团体内配对的结果被团体内所有成员观察到。但是,一个团体无法观察到另一团体的成员的博弈历史。这就使得团体间的合作很难维持,因为那些在团体间配对中有过背叛行为的人,无法被挑出来接受对方团体中成员的惩罚。但是, Fearson 和 Laitin 指出,即使没有这种直接惩罚,团体间合作依然有可能维持下去。他们分析了这一博弈的两个这种均衡。第一个是螺旋均衡。在这一均衡中,在团体内部,通过对背叛者施加 k^{in} 个时期的惩罚,实现团体内合作。在团体之间,通过实施 k^{out} 个时期的针对背叛团体的惩罚,实现团体之间的合作。如果有人在团体间配对博弈中发生背叛,则在惩罚时期中,背叛者所在团体中的所有成员被另一团体惩罚。另一种均衡是团体内约束均衡。在这个均衡中,不存在团体间惩罚,但是每个团体都就自己成员的背叛行为惩罚自己。我们只分析团体内约束均衡,而建议读者阅读 Fearon 和 Laitin 的原始论文以了解他们对螺旋均衡的讨论。

9.6.1 团体内约束均衡

我们只分析团体 A 的行为,关于 A 的结论很容易推广至团体 B。 $s_t = (k_1, k_2, \dots, k_n)$ 表示系统状态, k_i 是在时期 t 开始时参与者 i 的惩罚期中剩余的时期数。如果 $k_i = 0$, 参与者 i 是个合作者; 如果 $k_i > 0$, 参与者 i 是个背叛者。团体内约束均衡中的策略由良好行为构成(就是说,没有人主动背叛)。在偏离均衡的路径上,背叛者在某一数量的时期中,受到自己团体内的成员的惩罚。背叛者与惩罚他们的合作者保持合作; 实施惩罚的人不与背叛者合作。所有类型的成员都与团体外的成员合作。Fearon 和 Laitin 是这样描述这一策略的:

在团体之间的所有的配对中使用策略 C。在团体内部的配对博弈中,对不在惩罚阶段的成员使用策略 C,对在惩罚阶段的成员使用策略 D。一个成员,如果背叛了团体外的合作者或团体内的合作者,就进入或重新开始持续 k^{out} 个时期的惩罚阶段。

给定状态 s_t 和任意整数 $l > 0$, 设 n_{t+l} 是团体 A 中在时期 $t+l$ 里为合作者的成员数量,假设每个人从时期 t 到时期 $t+l$ 都使用均衡策略。于是, $q_{t+l} = \frac{n_{t+l}}{n-1}$ 是在

团体内配对博弈中面对着的对手是合作者的概率。

这些策略构成子博弈精炼 Nash 均衡。要证明这一点,我们必须验证以下条件:

1. 合作者 i 没有动机
 - (a) 背叛任何一位团体外部的参与者,
 - (b) 背叛任何一位团体内部的合作者,
 - (c) 与任何一位团体内部的背叛者合作。
2. 背叛者 i 没有动机
 - (a) 背叛任何一位团体外部的参与者,
 - (b) 背叛任何一位团体内部的合作者,
 - (c) 与任何一位团体内部的背叛者合作。

条件 1(c) 和条件 2(c) 显然得到满足;条件 1(c) 和条件 2(c) 中的偏离行为降低当前时期中的效用但是不影响其他参与者的策略,就是说,不致受到惩罚。条件 1(a) 和条件 1(b) 反映了相同的替代关系;条件 1(a) 和条件 1(b) 中的背叛行为在时期 t 产生收益 α ,接着便进入持续 k^{sp} 个时期的惩罚期。所以,我们只需证明在条件 1(a)、条件 2(a) 和条件 2(b) 三种情况中,没有偏离动机。

与团体内的合作者或外团体的成员的合作,带来的效用为:

$$1 + \sum_{l=1}^{\infty} \delta^l (p + (1-p))(q_{t+l} + (1-q_{t+l})\alpha)$$

而从条件 1(a) 和条件 1(b) 中的偏离行为所实现的效用为:

$$\begin{aligned} \alpha + \sum_{l=1}^{k^{sp}} \delta^l (p + (1-p)(-q_{t+l}\beta + (1-q_{t+l})0)) \\ + \sum_{l=k^{sp}+1}^{\infty} \delta^l (p + (1-p))(q_{t+l} + (1-q_{t+l})\alpha) \end{aligned} \quad (9.1)$$

合作所实现的净效用因此为:

$$1 - \alpha + \sum_{l=1}^{k^{sp}} \delta^l ((1-p)(q_{t+l}(1+\beta) + (1-q_{t+l})\alpha)) \quad (9.2)$$

所以,如果条件 1(a) 和条件 1(b) 中的偏离无利可图,则对所有的状态和所对应的数列 q_{t+l} , (9.2) 式肯定为正。如果 $\alpha > 1 + \beta$, 则在 $q_{t+l} = 1$ 时, (9.2) 式被最小化,其中 $l = 1, \dots, k^{sp}$ 。这就是自 $s_t = (0, 0, \dots, 0)$ 之后的路径。当且仅当:

$$\delta^{k^{sp}} \leq 1 - \frac{(1-\delta)(\alpha-1)}{\delta(1-p)(1+\beta)} \quad (9.3)$$

在这一状态之后净效用为正。设 $1 + \beta > \alpha$ 。在 $q_{t+l} = 0$ 时,净效用被最小化,其中 $l = 1, \dots, k^{sp}$ 。根据 q 的定义,这是个不可行的数列,因为所有参与者都被假定在

其惩罚阶段中选择合作,在 k^{sp} 个时期后结束自己的惩罚期。因此,在这个最小化净效用的数列中,所有的参与者都在时期 $t-1$ 背叛,在时期 $t+k^{sp}-1$ 恢复合作。因此,在数列 q 中,在 $l=1$ 时和 $l=k^{sp}-1$ 时, $q_{t+l}=0$, $q_{t+k^{sp}}=1$ 。进行代数运算,得到以下条件:

$$\frac{\delta - \delta^{k^{sp}}}{1 - \delta} \frac{\alpha}{1 + \beta} + \delta^{k^{sp}} \geq \frac{\alpha - 1}{(1 - p)(1 + \beta)} \quad (9.4)$$

现在分析在时期 t ,背叛者是否想进行条件 2(a) 中的偏离。假设处于状态 k_i 的背叛者与团体外的参与者配对。合作带给他的效用是:

$$1 + \sum_{l=1}^{k_i-1} \delta^l (p + (1-p)(-q_{t+l}\beta + (1-q_{t+l})0)) \\ + \sum_{l=k_i}^{\infty} \delta^l (p + (1-p)(q_{t+l} + (1-q_{t+l})\alpha))$$

偏离带给他的效用由 (9.1) 式给出。因此,只要:

$$\sum_{l=k_i}^{k^{sp}} \delta^l ((1-p)(q_{t+l}(1+\beta) + (1-q_{t+l})\alpha)) \geq \alpha - 1$$

背叛者就会与外团体成员合作。在 $k_i = k^{sp}$ 时,这一不等式的左边部分被最小化,因此:

$$\delta^{k^{sp}} ((1-p)(q_{t+k^{sp}}(1+\beta) + (1-q_{t+k^{sp}})\alpha)) \geq \alpha - 1$$

由于所有参与者都使用均衡策略,所以对所有的 s_i , $q_{t+k^{sp}} = 1$ 。因此,除非:

$$\delta^{k^{sp}} \geq \frac{(\alpha - 1)}{(1 - p)(1 + \beta)} \quad (9.5)$$

否则,背叛就有利可图。

最后,我们还需要排除掉条件 2(b) 中的偏离行为。假设有状态为 $k_i = 1$ 的背叛者与一位合作者配对博弈。合作实现的效用为:

$$-\beta + \sum_{l=1}^{k_i-1} \delta^l (p + (1-p)(-q_{t+l}\beta + (1-q_{t+l})0)) \\ + \sum_{l=k_i}^{\infty} \delta^l (p + (1-p)(q_{t+l} + (1-q_{t+l})\alpha))$$

背叛实现的效用为:

$$\sum_{l=1}^{k^{sp}} \delta^l (p + (1-p)(-q_{t+l}\beta + (1-q_{t+l})0)) \\ + \sum_{l=k^{sp}+1}^{\infty} \delta^l (p + (1-p)(q_{t+l} + (1-q_{t+l})\alpha))$$

合作实现的净效用因此为:

$$\sum_{t=k_1}^{k^{sp}} \delta^t ((1-p)(q_{t+1}(1+\beta) + (1-q_{t+1})\alpha)) - \beta$$

根据条件 2(a) 中的分析思路, SPNE 要求:

$$\delta^{k^{sp}} \geq \frac{\beta}{(1-p)(1+\beta)} \quad (9.6)$$

我们得到了全部的均衡条件。首先, 在 $\alpha > 1 + \beta$ 时, SPNE 要求 (9.3) 式、(9.5) 式和 (9.6) 式成立。只要 (9.5) 式成立, (9.6) 式就成立。而且, 我们还能证明, (9.5) 式意味着 (9.3) 式成立。我们可以把 (9.3) 式表述为:

$$\delta \left(\delta^{k^{sp}} - \frac{(\alpha-1)}{(1-p)(1+\beta)} \right) \leq \delta - \frac{(\alpha-1)}{(1-p)(1+\beta)} \quad (9.7)$$

如果 (9.5) 式成立, 则 (9.7) 式中不等号两边都为正; 由于 $1 > \delta > \delta^{k^{sp}}$, 所以, 右边部分肯定更大。这就是说, 在 $\alpha > 1 + \beta$ 时, (9.5) 式是团体内约束策略构成子博弈精炼 Nash 均衡的充分必要条件。

在 $\alpha < 1 + \beta$ 时, SPNE 要求 (9.4) 式、(9.5) 式和 (9.6) 式成立。(9.6) 式意味着 (9.5) 式成立。并且, 如果 (9.6) 式成立, 则下面的不等式成立:

$$\delta^{k^{sp}} \geq \frac{\beta}{(1-p)(1+\beta)} > \frac{\alpha-1}{(1-p)(1+\beta)} > \frac{\alpha-1}{(1-p)(1+\beta)} - \frac{(\delta - \delta^{k^{sp}})}{1-\delta} \frac{\alpha}{(1+\beta)}$$

就是说, (9.6) 式蕴含 (9.4) 式, 所以, (9.6) 式是团体内惩罚构成子博弈精炼均衡的充分必要条件。我们实际上证明了下面的命题:

命题 9.1 惩罚期为 k^{sp} 个时期的团体内惩罚策略构成子博弈精炼 Nash 均衡, 当且仅当:

$$\delta^{k^{sp}} \geq \min \left\{ \frac{\alpha-1}{(1-p)(1+\beta)}, \frac{\beta}{(1-p)(1+\beta)} \right\}$$

这个 SPNE 有几个特征值得指出。首先, 在 $k^{sp} > 1$ 时, 如果有不等式:

$$\delta^{k^{sp}} \geq \min \left\{ \frac{\alpha-1}{(1-p)(1+\beta)}, \frac{\beta}{(1-p)(1+\beta)} \right\}$$

则在 $k^{sp} = 1$ 时, 这一不等式也肯定成立。就是说, 要实现这一均衡, 只需要惩罚一个时期就足矣。事实上, 更长的惩罚期是反生产性的, 会降低背叛者为结束惩罚而选择合作的动力。

这一子博弈精炼 Nash 均衡的另一个重要特征是, 只有在与外部团体之间进行博弈的概率 p 足够小时, 这一均衡才能实现。由于这一 SPNE 要求:

$$\min \left\{ \frac{\alpha-1}{(1-p)(1+\beta)}, \frac{\beta}{(1-p)(1+\beta)} \right\} \leq 1$$

所以,如果:

$$p > \min \left\{ \frac{1}{(1+\beta)}, 1 - \frac{\alpha-1}{(1+\beta)} \right\}$$

这一均衡不存在。

当与外部团体之间进行博弈的概率很大时,因为惩罚只由团体内的成员实施,所以对偏离行为实施惩罚的概率就很小。这一效应产生了与直觉有些相悖的一个结论:太多的团体之间的合作阻碍团体之间的合作。Fearon 和 Laitin 认为,这一结论对种族保持自己的种族边界的愿望提供了内生理由。

9.7 应用:贸易战

表 9.6 是对自由贸易博弈的推广。

表 9.6 自由贸易博弈

1\2	自由贸易	保 护
自由贸易	Θ_1, Θ_2	0, $\Theta_2 + \rho$
保 护	$\Theta_1 + \rho, 0$	ρ, ρ

其中, Θ_i 是对方开放市场对国家 i 具有的价值, ρ 是每个国家从保护自己的市场中获得的收益。和前面一样,触发策略能够实现自由贸易,当且仅当对这两个国家:

$$\frac{\Theta_i}{1-\delta} \geq \Theta_i + \rho + \frac{\delta\rho}{1-\delta} \text{ 或者 } \delta \geq \frac{\rho}{\Theta_i}$$

要实现这一均衡,关键在于每个国家能完全观察到对手的政策。这不是个很现实的假设,因为各个国家可能会使用看不见的贸易壁垒。同时,由于贸易流量随许多与贸易政策无关的市场因素而波动,所以,两个国家之间贸易量的下降未必是因为对方做了手脚。

为分析这些问题,假设每个国家无法直接观察到对方的政策,而只能观察到自己的贸易值 Θ_i 。设 Θ_i 是个随机变量。为使分析尽可能地简单,假设在国家 j 实行自由贸易时, $\Theta_i = \theta > 0$ 的概率为 π ; $\Theta_i = 0$ 的概率为 $1 - \pi$ 。如果国家 j 保护自己的市场, $\Theta_i = 0$ 的概率为 1。因此,如果观察到 $\Theta_i = \theta$, 国家 i 就能推断国家 j 在推行自由贸易;如果观察到 $\Theta_i = 0$, 国家 i 并不确定国家 j 的政策选择。设 $\pi\theta > \rho$ 。就是说,每个国家都偏好自由贸易的结果而不是双边保护主义政策。^①

显然存在一些 SPNE, 其中,每个国家在每个时期都实行贸易保护政策。我们来分析是否存在实现某种程度的自由贸易的均衡。这种均衡要求在观察到 $\Theta_i = 0$ 时,实行某种形式的惩罚。

首先看触发策略:国家 i 在观察到 $\Theta_i = 0$ 后,永久地推行贸易保护政策。在这

① 这个模型建立在 Green 和 Porter(1984)关于经济卡特尔的不完全串谋和价格战的基础上。

一策略下,第一个时期中的自由贸易带来的收益是 $\pi\theta + (1-\pi)0$ 。只要 $\Theta_1 = \Theta_2 = \theta$, 自由贸易就能持续到下一个时期,而事件 $\Theta_1 = \Theta_2 = \theta$ 的发生概率是 π^2 。因此,国家 i 在第二个时期中实现的收益为:

$$\pi^2(\pi\theta + (1-\pi)0) + (1-\pi^2)\rho = \pi^3\theta + (1-\pi^2)\rho$$

将这一逻辑应用于第三个时期中,得到:

$$\pi^5\theta + (1-\pi^4)\rho$$

于是,自由贸易产生的效用的无限贴现和为:

$$V^{FT} = \pi\theta + \delta(\pi^3\theta + (1-\pi^2)\rho) + \delta^2(\pi^5\theta + (1-\pi^4)\rho) + \dots$$

调整后,得到:

$$V^{FT} = \pi\theta(1 + \delta\pi^2 + \delta^2\pi^4 + \dots) + \delta(1-\pi^2)\rho + \delta^2(1-\pi^4)\rho + \dots$$

等号右边的第一个级数等于:

$$\frac{\pi\theta}{1-\delta\pi^2}$$

第二个级数则可简化为:

$$\frac{\delta\rho}{1-\delta} - \frac{\delta\pi^2\rho}{1-\delta\pi^2}$$

因此,

$$V^{FT} = \frac{\pi\theta - \delta\pi^2\rho}{1-\delta\pi^2} + \frac{\delta\rho}{1-\delta}$$

偏离自由贸易转而实行保护政策所实现的效用更容易算出。发生偏离的时期实现的收益是 $\pi\theta + \rho$, 未来收益是 $\frac{\delta\rho}{1-\delta}$, 所以:

$$V^P = \pi\theta + \frac{\rho}{1-\delta}$$

因此,当且仅当 $V^{FT} \geq V^P$ 或者:

$$\delta > \frac{\rho}{\pi^3\theta}$$

国家 i 会选择自由贸易。

如果双方的政策能够被观察到,只要 $\delta > \frac{\rho}{\pi^3\theta}$, 就存在使用触发策略的子博弈精炼 Nash 均衡。所以,在政策无法被观察到时,触发策略均衡更难实现。事实上,如果 $\rho > \pi^3\theta$, 就不存在任何贴现率能使触发策略构成 SPNE。并且,触发策略的成本很高,因为即使由于偶然因素而导致很差的 Θ_i , 都会诱发无限期惩罚。即使没有人偏离自由贸易政策,这一均衡也几乎肯定会导致贸易

停止。

遵循 Green 和 Porter(1984), 我们探讨有限期触发策略。如果国家 i 观察到 $\Theta_i = 0$, 贸易战随之开始——在 $k \geq 1$ 个时期里, 两个国家都保护自己的市场。我们现在确定在什么条件下, 如果没有永久性贸易战, 自由贸易依然是最优政策。^① 设 V_i^k 是在持续 k 个时期的贸易战中, 国家 i 得到的这 k 个时期里的收益值。容易算出:

$$V_i^k = \frac{(1-\delta^k)\rho}{1-\delta}$$

设 V^{FT} 是两个国家处在自由贸易状态时, 自某个时期开始, 所实现的无限收益流的期望值。于是:

$$V^{FT} = \pi(\theta + \pi\delta V^{FT}) + (1-\pi^2)\delta(V_i^k + \delta^k V^{FT})$$

通过简单的代数运算, 得到:

$$V^{FT} = \frac{\pi\theta + (1-\pi^2)\delta V_i^k}{1-\pi^2\delta - (1-\pi^2)\delta^{k+1}}$$

根据单偏离原则, 偏离所实现的价值 V^P 为:

$$V^P = \pi\theta + \rho + \delta(V_i^k + \delta^k V^{FT})$$

在均衡中, $V^{FT} \geq V^P$ 。因此:

$$\frac{\delta - \delta^{k+1}}{1 - \delta^{k+1}} > \frac{\rho}{\pi^3\theta} \quad (9.8)$$

这一不等式的左边在 k 上递增, 所以设 $k^{\min}(\delta)$ 是对给定的 δ , 使这一不等式成立的最小的整数 k 。我们实际上已经证明了下面的命题:

命题 9.2 如果 $\delta > \frac{\rho}{\pi^3\theta}$ 且 $k \geq k^{\min}(\delta)$, 以下策略构成 SPNE:

- (1) 博弈开始时, 都实行自由贸易政策;
- (2) 两个国家都实行自由贸易, 直至对其中一个国家 i , $\Theta_i = 0$;
- (3) 在 $\Theta_i = 0$ 的时期之后的 k 个时期中, 实行保护政策;
- (4) 在 k 个时期后, 恢复自由贸易政策。

这里要注意的是, 只要贸易战的持续时间大于 $k^{\min}(\delta)$, 就能够实现 SPNE。但是, 只要两个国家能够协调到最优时长的贸易战上, 这一模型就提供了关于贸易战持续时间的理论。直观地说, 由于贸易战需要付出代价, 所以两个国家应该协调到能够实现合作的最短时长的贸易战上, 即 $k^{\min}(\delta)$ 上。这一点能够证明, 因为只要 $\pi\theta > \rho$, V^{FT} 就在 k 上严格递减。因此, 通过分析(9.8)式, 我们能得到经验实证结

① 由于两个国家都推行贸易保护也构成 Nash 均衡, 所以我们不需分析在贸易战中保护政策的最优性。

论。前面说过, (9.8) 式左边部分在 k 上递增; 这就意味着 $k^{\text{min}}(\delta)$ 在保护主义的值上递增, 在自由贸易的值上递减。这一结论合乎逻辑; 在激励问题更加严重时, 贸易战的时间长度就应该更长。贸易冲突的时长在 π 上递减。如果把 $1 - \pi$ 理解为贸易流量的波动性(在自由贸易体制下, 低贸易流量的发生概率), 则根据这种解释, 贸易波动性增加贸易战的时长——要阻止各个国家建立贸易壁垒和把贸易量的下降归诿过于自然波动, 就必须这样做。在均衡中, 只要不处于贸易战状态中, 两个国家都不推行保护政策, 所以, 他们知道, $\Theta_i = 0$ 是由自然波动导致的。有意思的是, 在均衡中, 虽然国家 j 从不背叛, 但是对 $\Theta_i = 0$, 国家 i 必须以高成本的和持续一定时期的贸易战方式做出反应; 要保证各个国家不设置贸易壁垒, 就必须这样做。

9.8 习题

习题 9.1 设无限重复博弈由如下的阶段博弈构成:

假设这两个参与者有相同的贴现因子

$\delta (0 < \delta < 1)$ 。

(1) 在这一重复博弈中是否存在这样的 Nash 均衡: 在每个时期中, 都使用策略组合 (u, m) ? 答案决定于贴现因子 δ 的值吗?

1\2	l	m	r
u	1, 1	8, 3	1, 2
c	1, 6	5, 5	1, 4
d	1, 2	4, 1	2, 2

(2) 给定如下策略: 在第一个时期中使用策略组合 (c, m) , 在任何时期中, 如果此前时期使用的是 (c, m) , 则使用 (c, m) 。如果在此前时期中, 没有使用 (c, m) , 则在以后所有时期中使用 (d, r) 。要使这个策略构成这一重复博弈的 Nash 均衡, δ 至少必须为多大?

(3) 在 Nash 均衡中, 如果每个时期里使用的策略组合都是 (u, m) , 则 δ 至少必须为多大?

习题 9.2 证明: 在表 9.4 的囚徒两难博弈中, 对以牙还牙策略的最优偏离导致以下两个收益流中的一个: $c + \delta d + \frac{c + \delta d}{1 - \delta^2}$ 和 $c + \frac{\delta b}{1 - \delta}$ 。

习题 9.3(*) 证明: 在无限重复博弈中, 如果参与者使用的贴现率小于 1, 则单偏离原则适用于子博弈精炼均衡。

习题 9.4 在重复的囚徒两难博弈中, 假设收益由表 9.4 给出。要使以牙还牙策略构成子博弈精炼均衡, 则贴现因子的取值范围是什么?

习题 9.5 证明命题 9.1。

习题 9.6 证明命题 9.2。

习题 9.7 在 Fearson 和 Laitin 模型中, 找出螺旋 SPNE 的存在性条件。这一均衡建立在如下策略基础上:

在团体内配对博弈中, 不管自己是合作者还是背叛者, 面对着合作者, 总是使用 **C**; 面对着背叛者, 总是使用 **D**。如果背叛合作者, 就进入或重新开

始为期 k^{in} 个时期的团体内惩罚阶段。在与外团体成员的配对博弈中,如果两个团体都不处于来自外团体的惩罚阶段中,就使用 C ; 否则,就使用 D 。在两个团体都不处于来自外团体的惩罚阶段中时,一个团体,如果有成员在团体间配对博弈中发生背叛,就进入到为期 k^{out} 个时期的来自外团体的惩罚阶段中。

习题 9.8 在贸易战模型中,构造如下“混合触发策略 SPNE”:在第一次观察到 $\Theta_i = 0$ 时,国家 i 没有永久地推行保护主义,而是使用混合策略,以 μ 的概率永久推行保护主义。



▶ 10

讨价还价理论

如果说政治科学研究“谁在何时以何种方式得到了什么”，讨价还价理论就是其基础。^①议员和政府官员就新法案讨价还价；各个国家为达成新的国际协定和解决危机讨价还价；各个政党为联合执政讨价还价；如此等等。

由于讨价还价的重要性，应用关于讨价还价的博弈论模型研究政治过程，成为了十分活跃的研究领域。这些模型探讨两类问题。第一类是分配问题——“谁胜”和“谁输”。总统得到了自己想要的法案吗？哪个国家控制冲突区域？哪些政党瓜分了政府职位？第二类问题则与政治上的讨价还价的效率有关。讨价还价过程本身消耗资源？或者没有达到改善每个人的福利的结果吗？即使存在着各方都偏好的妥协方案，立法上的讨价还价依然走入了死胡同或者被否决了吗？国际冲突最终演变为高成本的军事冲突和战争了吗？为什么要用如此长的时间才能组成新的联合政府？

本章介绍一些最重要的讨价还价模型和它们在政治科学中的应用。

10.1 Nash 讨价还价解

John Nash 建立的框架分析是最早的讨价还价模型。他的方法是公理性的，规定了任何讨价还价过程所产生的结果应该具有的一些特征。在讨论这些公理条件前，我们介绍 Nash 给出的讨价还价问题的解。我们的讨论仿照 Muthoo(1999)。

设两个参与者 A 和 B 就 X 数量的某种资源的分配在谈判。 X 可被无限细分；可行的分配是 x_A 和 x_B 且 $x_A + x_B \leq X$ 。每个人的所得决定每个人的效用，因此，两个人的效用分别为 $u_A(x_A)$ 和 $u_B(x_B)$ 。对参与者 $i \in \{A, B\}$ ，效用函数 $u_i(\cdot)$ 严格递增且是凹的。如果没有达成一致，每个人得到缺省效用——即分歧值或外部机会 $\underline{u}_i > u_i(0)$ 。为保证这一讨价还价问题有意义，设至少存在一种分配方案 (x_A, x_B) ，使得对每个 i ， $u_i(x_i) > \underline{u}_i$ 且 $x_A + x_B \leq X$ 。就是说，相对于各自的分歧值，至少存在一个可行分配方案被每个人所偏好。

^① 见 Lasswell(1936)。

为探讨 Nash 解,我们可以将这一问题转化为对效用的分配,即 (u_A, u_B) ,而不是对 X 的分配。我们把可行的效用分配集定义为:

$$\Omega = \{(u_A, u_B): \text{存在}(x_A, x_B) \text{使得 } u_A(x_A) = u_A, u_B(x_B) = u_B, x_A + x_B \leq X\}$$

根据对效用函数的假定,这一可行集的边界是如图 10.1 中的轨迹,也就是函数 $g(u_A) = u_B(X - u_A^{-1}(u_A))$ 。Muthoo(1999)证明了 g 在 u_A 上是递减和凹的。为简化分析,假定 g 二阶可导。

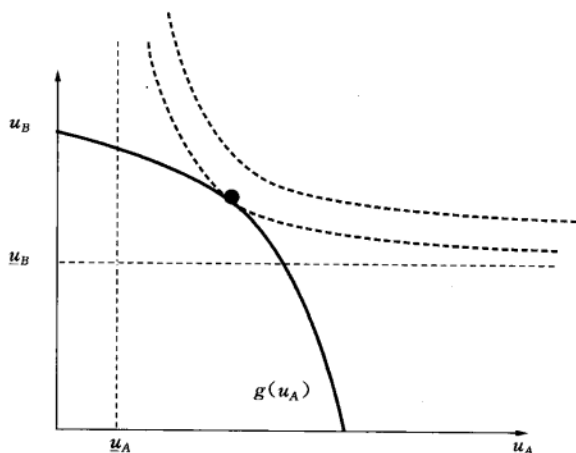


图 10.1 Nash 的讨价还价解

我们现在介绍 Nash 给出的这一讨价还价问题的解。它建立在稍后将讨论的几个公理的基础上,是在 $u_A \geq \underline{u}_A$ 和 $u_B \geq \underline{u}_B$ 的约束下,最大化 $(u_A - \underline{u}_A)(u_B - \underline{u}_B)$ 的效用分配方案 $(u_A, u_B) \in \Omega$ 。在图 10.1 中,条件 $u_A \geq \underline{u}_A$ 和 $u_B \geq \underline{u}_B$ 用虚线表示。因此,约束集是区间 $[\underline{u}_A, g^{-1}(\underline{u}_B)]$ 上的 $g(u_A)$ 。由于 g 是凹的和递减的,所以这一可行集是凸的。容易证明, Nash 积 $(u_A - \underline{u}_A)(u_B - \underline{u}_B)$ 在 u_A 和 u_B 上是拟凹的。因此,在这个积取正值时,它的等高线是图 10.1 中的粗虚线。

唯一的 Nash 讨价还价解位于 g 和等积曲线的切点上。在数学上,这一约束条件下的优化问题的解为:

$$-g'(u_A) = \frac{u_B - \underline{u}_B}{u_A - \underline{u}_A}$$

$$u_B = g(u_A)$$

在讨论 Nash 讨价还价解的一般结论前,我们看一种特殊情况。设 $X = 1$ 且 $u_i(x_i) = x_i$ 。Nash 讨价还价解为:

$$u_A = x_A = \frac{1 + \underline{u}_A - \underline{u}_B}{2} \quad \text{和} \quad u_B = x_B = \frac{1 - \underline{u}_A + \underline{u}_B}{2}$$

这个解有两个重要特征。首先,每个参与者,在分歧为其提供的效用越高时,在这

一解中实现的福利就越高;在对手有更好的外部机会时,所实现的福利就越差。其次,如果双方有相同价值的外部机会,资源就在两人之间平分。这里,双方以自己的分歧值为底线,平分剩余部分,即平分 $1 - \underline{u}_A - \underline{u}_B$ 。这种分配带给双方的效用为

$$\underline{u}_i + \frac{1 - \underline{u}_A - \underline{u}_B}{2}。$$

回到一般情况中。下面的定理给出了(用份额表示的)Nash 讨价还价解:

定理 10.1 Nash 讨价还价解是方程:

$$\frac{u_A(x_A) - \underline{u}_A}{u'_A(x_A)} = \frac{u_B(X - x_A) - \underline{u}_B}{u'_B(X - x_A)}$$

的解。

证明:根据前述最大化问题的解,可以得到:

$$g'(u_A) = -u_A^{-1}(X - u_A^{-1}(u_A)) \cdot u'_B(X - u_A^{-1}(u_A)) = \frac{-u'_B(X - u_A^{-1}(u_A))}{u'_A(X - u_A^{-1}(u_A))}$$

$$u_A^{-1}(u_A) = x_A$$

利用这一事实可证明这一命题。□

这一定理的一个直接结论是,如果这两个参与者有相同的分歧值和效用函数, Nash 讨价还价解为 $x_A = x_B = \frac{1}{2}X$ 。最后,根据我们关于 g 的假设,每个人的收益在自己的分歧值上递增,在对手的分歧值上递减。

定理 10.2 假设 g 二阶可导,则 $\frac{\partial u_i}{\partial u_i} > 0$; 如果 $i \neq j$, 则 $\frac{\partial u_i}{\partial u_j} < 0$ 。

证明:由于 g 二阶可导,我们可以对解:

$$\frac{g(u_A) - \underline{u}_B}{u_A - \underline{u}_A} + g'(u_A) = 0$$

进行隐函数求导。由于二阶条件得到满足,就是说:

$$\frac{g'(u_A)(u_A - \underline{u}_A) - (g(u_A) - \underline{u}_B)}{(u_A - \underline{u}_A)^2} + g''(u_A) < 0$$

于是,从:

$$\frac{-1}{u_A - \underline{u}_A} < 0, \quad \frac{(g(u_A) - \underline{u}_B)}{(u_A - \underline{u}_A)^2} > 0 \quad \text{和} \quad g'(u_A) < 0$$

得到上述结论。□

10.1.1 应用:厌恶风险和 Nash 讨价还价解

风险是讨价还价的重要构成因素。讨价还价者始终必须面对着这样一种可能

性:双方无法达成一致,而只得到外部机会。某个参与者如果狮子大开口,就会增加谈判破裂的概率。所以,更愿意承受风险的参与者应该做得更好,因为他们的要求更苛刻并且拒绝的方案更多。Nash 讨价还价模型尽管把讨价还价过程作为“黑箱”对待,但是,其解与我们的直观认识相一致。

假设每个参与者有效用函数 $u_i(x_i) = x_i^\alpha$, 其中, $0 < \alpha_i < 1$, 有分歧值 $\underline{u}_i = 0$, $X = 1$ 。 α 的不同取值,反映了参与者不同的风险厌恶态度; α 越小,风险厌恶程度越大。^①利用定理 10.1 中的公式,均衡解为:

$$x_A = \frac{\alpha_B}{\alpha_A + \alpha_B}, x_B = \frac{\alpha_A}{\alpha_A + \alpha_B}$$

在这个解中,每个人的份额随自己的风险厌恶系数递减,随对手的风险厌恶系数递增。这一结果也与前述直观认识相一致:风险承受能力足够大因而立场强硬(就是说,增加了无法达成一致的可能性)的一方,得到更大份额。

10.1.2 Nash 的公理

在本节中,我们给出构成 Nash 讨价还价解的基础的几个公理。不严格地说,这些公理概括了以下原则:

- (1) 讨价还价者最大化期望效用。
- (2) 讨价还价是有效率的。参与者分配了所有可供分配的资源,没有人的所得低于其分歧值。
- (3) 分配方案只决定于参与者的偏好和分歧值。
- (4) 对 Nash 解以外的分配方案,可以被剔除出去,不予考虑。这不会影响讨价还价问题的解。

我们正式地表述这些公理。前面说过, Ω 是通过对 X 的某种分配能够实现的可行的效用水平 (u_A, u_B) 构成的集合。帕累托最优配置集为 $\Omega' = \{\omega \in \Omega: u_A \geq u_A, g(u_A) \geq \underline{u}_B\}$ 。一般的讨价还价问题可以表述为数对 $(\Omega, \underline{u})$, $\underline{u}$ 为分歧值向量。所有的讨价还价博弈构成集合 Σ 。讨价还价解为函数 $F: \Sigma \rightarrow \mathbb{R}^2$; F_i 表示参与者 i 得到的效用。

以下公理构成了 Nash 解的基础。

公理 10.1 效用函数的等价表示:定义 $u'_i = \alpha_i u_i + \beta_i$, $\underline{u}'_i = \alpha_i \underline{u}_i + \beta_i$, 其中 $\alpha_i > 0$; 据此相应地定义 Ω' 。则对 $i = A, B$, $F_i(\Omega', \underline{u}') = \alpha_i F_i(\Omega, \underline{u}) + \beta_i$ 。

效用函数和分歧值的仿射变换,不改变讨价还价的结果。由于效用分配是以相同的效用函数变换进行调整的,所以在两个讨价还价解中,资源分配相同。正式地说,

① 标准风险厌恶系数为 $-\frac{u''}{u'}$ 。对这里的效用函数, $-\frac{u''}{u'} = -\frac{\alpha(\alpha-1)x^{\alpha-2}}{\alpha x^{\alpha-1}} = \frac{1-\alpha}{x}$ 。

$$u_i^{-1}(F_i(\Omega, \underline{u})) = u_i'^{-1}F_i(\Omega', \underline{u}')$$

根据第3章的讨论,这一公理意味着参与者都是期望效用最大化者。

公理 10.2 帕累托效率 如果 $F(\Sigma) = (u_A, u_B)$, 则不存在其他配置 $(u'_A, u'_B) \in \Omega$ 使得对某个 i , $u'_i > u_i$, 对 $j \neq i$, $u'_j \geq u_j$ 。

帕累托公理指出,不存在其他的分配方案,能够在不降低某个人效用的同时,提高另一个人的效用。讨价还价者无法找到这样的方案和改善讨价还价解。

公理 10.3 对称性 设 $\underline{u}_A = \underline{u}_B$, $(u_1, u_2) \in \Omega$ 当且仅当 $(u_2, u_1) \in \Omega$, 则 $F_A(\Omega, \underline{u}) = F_B(\Omega, \underline{u})$ 。

这一公理的基本思想是,如果没有人拥有比对手更好的分歧值或效用水平,则他们得到相同的效用分配。

公理 10.4 无关方案的独立性 IIA: 设有两个讨价还价博弈 $(\Omega, \underline{u})$ 和 $(\Omega', \underline{u})$ 。如果 $\Omega' \subset \Omega$ 且 $F(\Omega, \underline{u}) \subset \Omega'$, 则 $F(\Omega, \underline{u}) \subset F(\Omega', \underline{u})$ 。

这一公理说的是,保持分歧点不变,更小的可行配置集,只当它使初始配置不可行时,才改变讨价还价解。

根据前一节中对 Nash 讨价还价解的分析,它显然满足上述所有公理。下面的定理指出,它是唯一的满足所有四个公理的解。

定理 10.3 讨价还价解 $F: \Sigma \rightarrow \mathbb{R}^2$ 满足公理 10.1~10.4 当且仅当它是 Nash 讨价还价解。

Muthoo(1999)给出了这一定理的直接证明。读者应该能够独立完成习题中所介绍的证明过程。

10.2 非合作讨价还价

从 Nash 讨价还价解中,能够得出许多合理的经验结论,但是,我们最好将其理解为一个规范命题,它告诉我们讨价还价结果应该具有的特征;而不是将其理解为实证理论,即关于实际的讨价还价过程如何进行的理论。在本节中,我们用非合作博弈论模型确定在不同的扩展式博弈中,讨价还价者的行为。

Rubinstein(1982)最早应用非合作博弈论分析讨价还价问题。两个参与者在决定如何分配 1 元钱。他们轮流提出分配方案,就是说,参与者 1 在时期 0, 2, 4, ... 中提出方案,参与者 2 在其他时期中提出方案。这一博弈持续进行(可能无限期地持续下去),直至某个分配方案被对方接受。

参与者 1 在由自己提出方案的每个时期中,提出方案 (x_1, x_2) , 其中, x_1 是参与者 1 的份额, x_2 是参与者 2 的份额, 且 $x_1 + x_2 \leq 1$ 。如果参与者 2 接受这一方案, 博弈结束, 这 1 元钱按照方案分配。参与者 2 如果拒绝这一方案, 则由他提出新的方案, 即新的 (x_1, x_2) 。参与者 1 如果拒绝这一方案, 博弈继续进行下去。为简化分析, 设两个人都有线性效用函数: $u_1(x_1, x_2) = x_1$ 和 $u_2(x_1, x_2) = x_2$ 。每个人有贴现因子 δ_i ; 在未来第 t 个时期提出的方案 (x_1, x_2) 带给参与者的效用为 $(\delta_1^t x_1, \delta_2^t x_2)$ 。

和第7章中的讨价还价博弈一样,在这一博弈中存在着大量 Nash 均衡。设一个策略组合为“参与者1要求 $x_1 = 1$ 并拒绝其他所有分配方案,参与者2提出 $x_1 = 1$ 并接受其他所有分配方案”。它构成 Nash 均衡,却不是子博弈精炼均衡。参与者2如果拒绝参与者1在首轮中提出的方案($x_1 = 1$),并提出 $x_1 > \delta_1$ 的分配方案,参与者1会接受这一方案;参与者2能得到的最好结果不过是下一个时期中的完整的1元钱乘以贴现因子 δ_1 。我们集中探讨子博弈精炼 Nash 均衡。

10.2.1 子博弈精炼均衡

Rubinstein 指出,在这一博弈中,存在唯一的 SPNE,它规定在每个时期中使用如下策略:

参与者1提出方案:

$$\left(\frac{1-\delta_2}{1-\delta_1\delta_2}, \frac{\delta_2(1-\delta_1)}{1-\delta_1\delta_2} \right)$$

并接受参与者2提出的方案当且仅当在参与者2的方案中:

$$x_1 \geq \frac{\delta_1(1-\delta_2)}{1-\delta_1\delta_2}$$

参与者2提出方案:

$$\left(\frac{\delta_1(1-\delta_2)}{1-\delta_1\delta_2}, \frac{1-\delta_1}{1-\delta_1\delta_2} \right)$$

并接受参与者1提出的方案当且仅当在参与者1提出的方案中:

$$x_2 \geq \frac{\delta_2(1-\delta_1)}{1-\delta_1\delta_2}$$

我们证明这种策略构成 SPNE。首先验证参与者1是否有动机在任何子博弈中偏离这一策略。设某个子博弈从参与者1提出方案开始(即从偶数时期开始)。在均衡中,参与者1提出的分配方案为:

$$\left(\frac{1-\delta_2}{1-\delta_1\delta_2}, \frac{\delta_2(1-\delta_1)}{1-\delta_1\delta_2} \right)$$

这一方案被参与者2接受。降低 x_1 显然无法给参与者1带来好处,参与者2会接受这一方案,而参与者1的所得低于均衡份额。参与者1如果提高 x_1 ,则他必须降低 x_2 才能使方案依然可行。但是,只要:

$$x_2 < \frac{\delta_2(1-\delta_1)}{1-\delta_1\delta_2}$$

方案就会被参与者2拒绝。参与者2在拒绝参与者1的方案后,提出 $x_1 = \frac{\delta_1(1-\delta_2)}{1-\delta_1\delta_2}$; 这一方案被参与者1接受。在这里,我们使用的是单偏离原则——改变方案提出

方的行为但不改变方案接受方的行为。因此,参与者 1 从自己的偏离行为中得到的效用是 $\frac{\delta_1^2(1-\delta_2)}{1-\delta_1\delta_2}$, 由于 $\delta_1 < 1$, 它小于均衡效用 $\frac{1-\delta_2}{1-\delta_1\delta_2}$ 。

再分析在参与者 2 首先提出方案时(即从奇数时期开始),参与者 1 是否会偏离上述策略。参与者 2 提出的分配方案为:

$$\left(\frac{\delta_1(1-\delta_2)}{1-\delta_1\delta_2}, \frac{1-\delta_1}{1-\delta_1\delta_2}\right)$$

此时,参与者 1 从接受或拒绝这一方案中得到相同的效用;对参与者 1 来说,最好的结果只是使下一时期中,方案 $x = \frac{(1-\delta_2)}{1-\delta_1\delta_2}$ 被对方接受。

证明参与者 2 不会偏离均衡策略的过程,与此完全相同。

10.2.2 计算均衡

前述证明没有告诉我们如何得到均衡结果。我们现在给出一个更富有构造性的证明。设在参与者 1 和参与者 2 分别为首轮方案提出者的子博弈中, v_1 和 v_2 是他们各自从相应子博弈中得到的效用;就是说,如果参与者 1 提出了被对方接受的方案 x_1 , 则 $v_1 = x_1$ 。参与者 1 的方案如果被拒绝,则在参与者 2 提出的 x_1 和参与者 1 拒绝并在下一轮中自己提出的 x_1 中,最大的那个值就是 v_1 。由于每个时期中所使用的策略相同,所以这些值与 t 无关。这些值是持续值,因为它们同时是拒绝对手方案从而进入下一个子博弈所实现的效用。我们分析参与者 1 是首轮方案提出者的子博弈。他提供给参与者 2 的效用至少得为 $\delta_2 v_2$ 。因此, $x_1 = 1 - \delta_2 v_2$ 。由于这一方案被接受,所以, $v_1 = x_1 = 1 - \delta_2 v_2$ 。再看参与者 2 是首轮方案提出者的子博弈。他提供给对方的效用至少得为 $\delta_1 v_1$, 因此, $v_2 = 1 - \delta_1 v_1$ 。

求解这两个方程,得到 $v_1 = \frac{1-\delta_2}{1-\delta_1\delta_2}$ 和 $v_2 = \frac{1-\delta_1}{1-\delta_1\delta_2}$ 。这些持续值与前一节中提出的策略相一致。事实上,前一节中提出的策略是获得这些持续值的唯一可行的方式。

10.2.3 唯一性

我们已证明 Rubinstein 均衡是子博弈精炼 Nash 均衡,但是,可能还存在其他子博弈精炼 Nash 均衡。我们现在通过证明 v_1 和 v_2 是与 SPNE 相一致的唯一的持续值,来证明 Rubinstein 均衡是唯一的 SPNE。假设存在不止一个 SPNE。设 $\bar{v}_i$ 和 $\underline{v}_i$ 是在参与者 i 为首轮方案提出者的任何子博弈中,参与者 i 所实现的最高的和最低的 SPNE 持续值。设 $\bar{w}_i$ 和 $\underline{w}_i$ 是在参与者 i 不是首轮方案提出者的任何子博弈中,参与者 i 所实现的最高的和最低的 SPNE 持续值。

在参与者 1 提出分配方案时,他提供给对方的效用不会高于 $\delta_2 \bar{v}_2$, 因为参与者 2

不可能通过拒绝参与者 1 的方案和在下一轮中提出自己的方案而得到高于 $\bar{v}_2$ 的效用。所以,参与者 1 的最低的持续值肯定满足 $\underline{v}_1 \geq 1 - \delta_2 \bar{v}_2$ 。同理, $\underline{v}_2 \geq 1 - \delta_1 \bar{v}_1$ 。由于对手提供给参与者 i 的效用水平永远不会超过 $\delta_i \bar{v}_i$, 所以 $\bar{w}_i \leq \delta_i \bar{v}_i$ 。

现在看参与者 1 的策略。在由他提出方案时,他能做到的最好的结果是给对手 $\delta_2 \underline{v}_2$ 或者是对方拒绝而得到 $\delta_1 \bar{w}_1$ 。因此,参与者 1 的持续值满足:

$$\bar{v}_1 \leq \max\{1 - \delta_2 \underline{v}_2, \delta_1 \bar{w}_1\} \leq \max\{1 - \delta_2 \underline{v}_2, \delta_1^2 \bar{v}_1\} = 1 - \delta_2 \underline{v}_2$$

同理, $\bar{v}_2 \leq 1 - \delta_1 \underline{v}_1$ 。所以,以下四个不等式必须得到满足:

$$\underline{v}_1 \geq 1 - \delta_2 \bar{v}_2$$

$$\underline{v}_2 \geq 1 - \delta_1 \bar{v}_1$$

$$\bar{v}_2 \leq 1 - \delta_1 \underline{v}_1$$

$$\bar{v}_1 \leq 1 - \delta_2 \underline{v}_2$$

第一个和第三个不等式相结合,得到 $\underline{v}_1 \geq 1 - \delta_2(1 - \delta_1 \underline{v}_1)$, 即 $\underline{v}_1 \geq \frac{1 - \delta_2}{1 - \delta_1 \delta_2}$ 。

第二个和第四个不等式相结合,得到 $\bar{v}_1 \leq 1 - \delta_2(1 - \delta_1 \bar{v}_1)$, 即 $\bar{v}_1 \leq \frac{1 - \delta_2}{1 - \delta_1 \delta_2}$ 。这

两个条件意味着 $\bar{v}_1 = \underline{v}_1 = \frac{1 - \delta_2}{1 - \delta_1 \delta_2}$ 。同样的方法能得到 $\bar{v}_2 = \underline{v}_2 = \frac{1 - \delta_1}{1 - \delta_1 \delta_2}$ 。这就是说,对每个参与者,存在唯一的持续值,上述策略构成唯一的 SPNE。

10.2.4 结论

这一模型给出了一个简单的博弈路径。在时期 0, 参与者 1 提出方案 $(\frac{1 - \delta_2}{1 - \delta_1 \delta_2}, \frac{\delta_2(1 - \delta_1)}{1 - \delta_1 \delta_2})$, 参与者 2 接受这一方案, 博弈结束。由于整个 1 元钱得到了分配, 并且在分配过程中, 不存在任何滞后, 所以, 这个子博弈精炼 Nash 均衡是有效率的。容易验证, 这一 SPNE 有以下特征:

(1) 两个参与者如果有相同的贴现因子, 则存在先发优势, 因为 $\frac{1 - \delta}{(1 - \delta)^2} > \frac{\delta(1 - \delta)}{1 - \delta}$ 。直观地说, 由于参与者 2 贴现未来, 参与者 1 提供给参与者 2 的份额只需等于参与者 2 在下个时期中作为方案提出者时能得到的份额。由于两个参与者完全相同, 所以, 参与者 2 得到的数量不及参与者 1。

(2) 两个参与者得到的份额在自己的贴现因子上递增, 在对手的贴现因子上递减; 越有耐性, 得到越多。在参与者 2 的贴现因子很高时, 参与者 1 必须给参与者 2 更高的份额才能使参与者 2 马上接受分配方案。而当参与者 1 的贴现因子很高时, 在由参与者 2 提出方案时, 参与者 2 必须给参与者 1 更高的份额才能达成一致。这就是说, 参与者 2 如果拒绝参与者 1 的方案, 只会降低自己得到的效用, 所

以,参与者 1 能够在第一个时期里得到更大份额。

(3) 如果 $\delta_1 = \delta_2 = \delta$, 则随着 δ 收敛于 1, 两个参与者得到的份额都收敛于 $\frac{1}{2}$ 。

随着两个参与者变得极具耐性, 如果在对手提出的分配方案中自己的份额小于下个时期中自己作为方案提出者能得到的份额, 他们就更不愿接受对手的方案。在极限点, 他们要求得到自己预期在下一个时期中能得到的份额。我们可以这样理解贴现因子收敛于 1 的情况: 提出方案和方案被拒绝后由对手提出方案的过程, 进行得很快, 使得拒绝方案的行为所导致的时间滞后效应为无穷小。在这种情况下, 均衡意味着平均分配, 与 Nash 讨价还价解相一致。

10.2.5 不对称的分歧值

在标准的 Rubinstein 博弈中, 在任何时期中, 双方如果没有达成一致, 得到的效用是零。我们在两个方面发展这一博弈模型。首先, 在达成一致之前的每个时期中, 双方得到配置 (d_1, d_2) , $d_1 + d_2 < 1$ 。当双方在分配方案 (x_1^*, x_2^*) 上达成一致后, 双方在此后以至无穷远的每个时期中, 都得到方案中规定的份额。与前一节中的模型不同, 在这里, 所分配的份额“马上被消费掉”。^①为使分析简单, 假定 $\delta_1 = \delta_2 = \delta$ 。于是, 在时期 t 达成分配协议带给双方的效用为:

$$\left(\frac{(1-\delta^{t-1})d_1 + \delta^t x_1^*}{1-\delta}, \frac{(1-\delta^{t-1})d_2 + \delta^t x_2^*}{1-\delta} \right) \quad ②$$

设 v_i 是在由 i 提出方案的时期中, i 得到的持续值。如果双方在这样一个时期中达成协议 (x_1^*, x_2^*) , $v_i = \frac{x_i^*}{1-\delta}$ 。再来分析参与者 2 接受或拒绝 x_2 的决策。参与者 2 如果接受 x_2 , 得到的效用为 $\frac{x_2}{1-\delta}$; 参与者 2 如果拒绝, 参与者 2 在当前时期中得到 d_2 , 在下个时期中得到持续值 v_2 。所以, 只要 $x_2 > (1-\delta)(d_2 + \delta v_2)$, 参与者 2 就接受 x_2 。再看参与者 1 的决策。如果他给参与者 2 的份额为可被对方接受的最低的份额, 即 $x_2 = (1-\delta)(d_2 + \delta v_2)$, 则自己的持续值就是:

$$v_1 = \frac{1 - (1-\delta)(d_2 + \delta v_2)}{1-\delta} = \frac{1}{1-\delta} - d_2 - \delta v_2$$

如果参与者 2 希望对方接受自己提出的方案, 就需要:

① 调整后的模型排除了这样的策略: 双方不断地在讨价还价, 寄希望于 d_i 的贴现和超过从所达成的单时期分配协议中得到的份额。我们能很容易地调整原始模型, 使其包含这样的假设: 所分配的是一个效用流量而不是一次性消费。我们只需使用原始模型, 同时假定参与者们得到的份额为 $\frac{1}{1-\delta}$ 。

② 我们假定, 任何一个协议导致的各个时期中的分配份额都相同。但是, 由于参与者风险中性, 所以, 双方可能接受产生相同的收益的随机分配方案。

$$v_2 = \frac{1}{1-\delta} - d_1 - \delta v_1$$

求解这两个方程,得到:

$$v_1 = \frac{1-d_2+\delta d_1}{1-\delta^2} = \frac{d_1}{1-\delta} + \frac{1-d_1-d_2}{1-\delta^2}$$

$$v_2 = \frac{1-d_1+\delta d_2}{1-\delta^2} = \frac{d_2}{1-\delta} + \frac{1-d_1-d_2}{1-\delta^2}$$

要证明这两个值确实是均衡持续值,我们就必须证明,每个参与者都偏好均衡分配方案而不是偏离这一方案和多得到一个时期的分歧值。所以,我们要求 $v_1 > d_1(1-\delta) + \delta^2 v_1$, 即 $v_1 > \frac{d_1}{1-\delta}$ 。同理,在均衡中, $v_2 > \frac{d_2}{1-\delta}$ 。这些事实很容易验证。我们可以用第 10.2.3 节中的技术来证明这是唯一的 SPNE。

这一均衡有许多性质类似 Nash 讨价还价解。每个参与者的持续值在自己的分歧值上递增,在对手的分歧值上递减。对这一均衡,我们还可以从剩余分配角度来理解。每个参与者的持续值有两个构成部分。第一部分是 $\frac{d_i}{1-\delta}$, 它是在没有达成一致时,每个人能得到的收入。第二部分是 $\frac{1-d_1-d_2}{1-\delta^2}$, 它是在参与者的外部机会值为 0 时,分配 $1-d_1-d_2$ 的博弈的均衡持续值。因此,对这一均衡的一种很有用的理解是,两个参与者得到他们有权得到的部分,而后就剩余部分讨价还价。

10.3 封闭规则下多数通过的讨价还价

Rubinstein 模型的一个重要特点是,要就分配方案达成一致,需要双方一致同意。这一特征把许多重要的政治现象排除在外,因为在许多政治活动中,要达成一致,只需简单多数或超多数通过的方案。Baron 和 Ferejohn(1989)将 Rubinstein 模型推广至有两个以上的讨价还价者的简单多数投票通过的情况中。

假设有 N (N 为奇数) 个参与者在讨价还价,任何分配方案必须得到至少 $n = \frac{N+1}{2}$ 票才能通过。Baron 和 Ferejohn 没有假定各参与者轮流提出方案,而是采用随机识别原则的讨价还价模式。根据这一模式,在每个时期中,每个参与者以相同概率 $\frac{1}{N}$ 被选中,提出自己的方案。在本节中,我们探讨封闭规则下的讨价还价,就是说,提案人在当前立法时期中提出一个或者接受或者不接受的方案。每个时期中提出的方案为 $(x_1, x_2, \dots, x_N)$, 其中, x_i 是参与者 i 得到的份额。可行性要求 $\sum x_i \leq 1$ 。这一方案如果被拒绝,则这一时期结束,贴现随即发生,在下一时期开始时挑选出新的提案人。后面将探讨开放规则下的讨价还价,它允许提案在

当前时期被修正。为简化分析,假定每个参与者有相同贴现因子 δ 。

这一博弈有许多子博弈精炼均衡。事实上,在 N 和 δ 很大时,任意分配方式都有一个 SPNE 提供支持。参与者如果有足够的耐性,他们就能设计出惩罚策略,使背叛者的份额为 0。但是,这种策略要求每个参与者知道这个博弈的整个(可能持续无限个时期的)历史,只有如此,他们才能知道哪些行动与所设计的惩罚策略相一致。遵循 Baron 和 Ferejohn 的分析,我们只探讨稳定均衡。在这一博弈的稳定均衡中:

(1) 提案人,不管博弈历史是什么,在每次由他提出方案时,始终提出相同的分配方案。

(2) 投票者只根据当前提案和对未来提案的预期进行投票。根据前一假设,未来每个时期中的提案有相同的结果分布。

这两个假设意味着博弈在每个时期都是从头开始。因此,每个参与者的持续值是这一博弈的期望效用。设 v_i 是参与者 i 的持续值。我们只探讨对称均衡,即对所有的 i , $v_i = v$ 。最后一点,我们只探讨投票者在投票阶段并不使用弱劣策略的均衡。就是说,投票者接受任何方案,只要方案带给他的效用至少和贴现持续值相同。换句话说,从方案中得到 $x_i \geq \delta v$ 的投票人,会投票支持这一方案;如果得到的份额小于 δv ,投票者就会投反对票。

给定上述投票策略,最优方案给提案人 $z = 1 - (n-1)\delta v$, 给另外 $n-1$ 个参与者每人 δv , 给剩下的每个人 0。假定提案人随机选择自己的联盟伙伴(这是个可证明的结论)。我们现在计算 v 。由于持续值等于下一时期开始的博弈的期望值,所以,它等于 z 乘以被选中为提案人的概率 $\frac{1}{N}$, 加上 δv 乘以被接纳到胜出联盟的概率 $\frac{n-1}{N}$, 加上 0 乘以剩余概率:

$$v = \frac{z}{N} + \frac{n-1}{N}\delta v$$

代入 z 并简化,得到:

$$v = \frac{1}{N}$$

因此,持续值是对这 1 元钱的等比例瓜分。由于 v 同时是这一博弈的期望效用,所以这一结论意味着讨价还价过程有效率,因为参与者效用之和达到最大。但是,在本章末的习题中我们指出,如果投票者厌恶风险,这一效率结论未必成立。

最后,根据所算得的 v ,我们可以计算出提案人得到的份额:

$$z = 1 - \delta \frac{n-1}{N} = 1 - \delta \frac{N-1}{2N}$$

要确保提案人提出的是可被接受的方案,我们必须验证 $z > \delta v$; 否则,提案人会把球踢到下一个时期。这一条件容易证实。

这个模型的一个重要意义在于它对提案权的预测。衡量提案权的一个指标是 z 和 δv 之间的差或 $1 - \delta \frac{N-1}{2N}$ 。首先要看到,提案权在 N 上递增。在 N 上升时,提案人就有更多的潜在的联盟伙伴可供他选择。这会提高加入胜出联盟的竞争程度,压低提案人必须让出的份额。其次,提案权在 δ 上递减。在 δ 上升时,投票者更愿意否决提案,等待亲自提出方案的时机。因此,提案人必须更加慷慨,只有如此,才可能使自己的方案通过。

10.3.1 超多数通过原则

我们可以发展这一模型,分析通过法案需要超多数投票通过的情况。假设提案要获得通过,支持票数必须为 $k > n$ 。容易看出,提案人现在得到的份额是:

$$z = 1 - (k-1)\delta v$$

持续值为:

$$v = \frac{z}{N} + \frac{k-1}{N}\delta v$$

可以算出:

$$v = \frac{1}{N}$$

由于超多数投票规则依然保持了多数通过投票博弈的对称性,所以,出现这种结果并不让人意外。但是,提案人的均衡份额现在下降到 $z = 1 - \frac{\delta(k-1)}{N}$ 。就是说,超多数通过的投票规则削弱了提案人的优势。

10.3.2 不对称的提案权

上述模型都假定所有立法者有相同的概率分配方案。在现实世界中,立法制度和在某些委员会和党派中的成员资格等因素,都影响某个立法者提出方案的概率。

为发展这些模型,假设所有参与者被分为 A 和 B 两个党派。党派 A 有 $N-m \geq \frac{n+1}{2}$ 位成员,是多数党。 A 的每个成员拥有的提案权为 $p > \frac{1}{N}$ 。党派 B 中有 m 个成员,每个成员拥有的提案权为 $q < \frac{1}{N}$ 。为保持一致性,设 $(N-m)p + mq = 1$ 。

我们依然假定对称性,就是说,有相同的概率提出方案的每个立法者,使用相同的策略,有相同的持续值。设两个党派的成员的持续值分别为 v_A 和 v_B 。假设 $v_A > v_B$ (稍后将证明 $v_A > v_B$)。给定这些持续值,党派 A 的成员投票支持带给自己的效用至少为 δv_A 的任何方案;党派 B 的成员投票支持带给自己的效用至少为 δv_B 。

的任何方案。给定这一策略并假定 $v_A > v_B$, 党派 A 的提案人给党派 B 的 m 成员每人 δv_B , 给党派 A 中的 $n-m-1$ 个成员每人 δv_A 。前面说过, $n = \frac{N+1}{2}$, 所以, 提案人得到:

$$z_A = 1 - (n-m-1)\delta v_A - m\delta v_B$$

党派 B 的提案人提供给党派 B 的 $m-1$ 个成员每人 δv_B , 给党派 A 的 $n-m$ 个成员每人 δv_A , 于是, 提案人自己得到:

$$z_B = 1 - (n-m)\delta v_A - (m-1)\delta v_B$$

显然, $z_A > z_B$ 。我们现在可以求出 v_A 和 v_B :

$$v_A = pz_A + p(n-m-1)\delta v_A + \frac{qm(n-m)\delta v_A}{N-m}$$

$$v_B = qz_B + (1-q)\delta v_B$$

于是, 我们有了四个方程和四个未知数。求解这一方程组相当容易, 但是, 解的形式却相当烦琐, 所以, 我们只看一个简单的例子。设 $N=3$, $m=1$ 。于是, $q=1-2p < \frac{1}{3}$ 。均衡条件为:

$$z_A = 1 - \delta v_B$$

$$z_B = 1 - \delta v_A$$

$$v_A = pz_A + \frac{q\delta v_A}{2}$$

$$v_B = qz_B + (1-q)\delta v_B$$

在一些烦琐的代数运算后, 得到:

$$v_A = \frac{(1-q)(1-\delta)}{2+q\delta-2\delta}$$

$$v_B = \frac{q(2-\delta)}{2+q\delta-2\delta}$$

我们来证明 $v_A \geq v_B$ 。当:

$$q < \frac{1-\delta}{3-2\delta} \leq \frac{1}{3} \text{ 时, } v_A \geq v_B$$

在 $\delta > 0$ 时, q 的上界值始终小于 $\frac{1}{3}$, 所以, 提案权上的不对称性必须很大, 只有如此, 才能使党派 A 占优势。原因是, 党派 A 的更大的提案权使得党派 A 的成员成为不受欢迎的联盟伙伴。因此, 成为提案人的概率必须足够大, 从而抵消这一效应。但是不难证明, v_A 在 q 上递减而 v_B 在 q 上递增。

要完成我们的分析,还需要探讨 $\frac{1-\delta}{3-2\delta} < q \leq \frac{1}{3}$ 时的情况。我们能够排除 $v_B > v_A$ 的可能性,因为 $v_B > v_A$ 意味着 B 的成员永远不会与提案人结盟。所以:

$$v_B = qz_B = q(1-\delta v_A)$$

$$v_A = pz_A + (1-p)\delta v_A = p(1-\delta v_A) + (1-p)\delta v_A$$

于是得到:

$$v_A = \frac{1-q}{2(1-\delta q)}, v_B = \frac{q(2-\delta-\delta q)}{2(1-\delta q)}$$

在这里,只有当 $q \geq \frac{1+2\delta}{3-2\delta} \geq \frac{1}{3}$ 时,才有 $v_B > v_A$ 。而这违背了关于 q 的初始假设。

因此,在 $\frac{1-\delta}{3-2\delta} < q \leq \frac{1}{3}$ 时,唯一可能的结果是 $v_A = v_B$ 。要支持这一均衡,党派 A 的提案人必须使用混合策略,即在与党派 A 中其他成员结成联盟和与党派 B 中成员结成联盟之间,随机做出选择。我们把这一均衡混合策略的计算留做习题。

10.3.3 不对称的否决权

立法制度中的一个制度现象是,某些参与者,如总统、上议院或法庭等,有权否决法案。在本节中,我们用一个简单例子说明如何把否决权加入到 Baron-Ferejohn 模型中。^①假设在三人立法机构中,有一位成员拥有绝对否决权,就是说,所有提案都得由他审批。设党派 B 拥有否决权。为使分析简单,我们再回到平等提案权情况中。

由于党派 B 拥有绝对否决权,所以任何提案人都得与党派 B 结盟,同时还须与党派 A 中至少一位成员结盟,所以:

$$z_A = 1 - \delta v_B$$

$$z_B = 1 - \delta v_A$$

计算持续值,得到:

$$v_A = \frac{1}{3}z_A + \delta \frac{1}{3}v_A$$

$$v_B = \frac{1}{3}z_B + \delta \frac{2}{3}v_B$$

于是:

$$v_A = \frac{3(1-\delta)}{\delta^2 - 9\delta + 9} \quad \text{和} \quad v_B = \frac{3-2\delta}{\delta^2 - 9\delta + 9}$$

在这里,只要 $\delta > 0$, 便有 $v_A < v_B$ 。

① Baron-Ferejohn 博弈的这一发展由 McCarty(2000a, 2000b)提出。

10.4 开放规则下的 Baron-Ferejohn 模型

在前面几节讨论的模型中,提案无法在当前立法时期得到修正。我们可以修改这一点,允许在最终投票之前修正提案。假设在每个方案提出后,从剩下的 $N-1$ 个立法者中随机选出一人。所选出的这位立法者有两个选择。他可以建议就提案进行最终投票表决;他可以提出新的提案或者修正案。修正案与原始提案同时交付投票表决。在这轮投票中胜出的提案将是下一个时期中提交议会投票表决的提案。在下一个时期中,又有一位立法者被随机选中,由他提出修正案或者建议投票。

提案者现在有两点必须考虑。首先,和前面一样,在最后一轮投票中,简单多数的立法人员必须从提案中得到其贴现持续值,只有这样,才能在最终投票中有简单多数的立法人员支持提案。其次,提案者必须精心设计提案以防止其他人提出修正案。要做到这一点,就必须使下一个提案人得到足够多的份额,从而使他愿意把原始提案提交表决而不是在下一个时期开始时提出自己的提案给立法机构。

为使分析简单,设 $N=3$ 。首先分析这样一种方案:提案人得到 z ,另外两位立法者每人得到 $\frac{1-z}{2}$,从而使每个立法者都支持这一提案。为求出最优的 z ,定义 $v_i^2(z)$ 为在新的时期开始时给参与者 i 的份额是 z 而给另两位立法者每人的份额是 $\frac{1-z}{2}$ 的提案所具有的持续值。我们只探讨对称均衡,所以可以忽略下标 i 。于是, $v^2(z)$ 是这一策略给第一个提案者的期望效用,和给成功地修正这一提案的任何人的期望效用。

根据这一定义,提案人必须给每个立法者至少 $\delta v^2(z)$ 的效用才能得到立法者的支持。否则,修正阶段中随机挑选出来的立法者会拒绝这一提案,而提出自己的提案,使自己得到 z 。因此,均衡要求 $\frac{1-z}{2} \geq \delta v^2(z)$ 。只要这一条件得到满足,提案人就能以 1 的概率得到 z ,从而使 $v^2(z) = z$ 。所以,提案人在条件 $\frac{1-z}{2} \geq \delta v^2(z)$ 约束下最大化 z 。这一最大化问题的解为:

$$v^2(z) = z = \frac{1}{1+2\delta}$$

在这里,提案者能确保自己得到 $z = \frac{1}{1+2\delta}$ 。他也可能只想取得其中一位立法者的支持,但是,这样做的风险是,如果被排除在联盟之外的立法者被随机选中而提出自己的修正案的话,初始提案也就落空。假设提案者自己得到 z ,另一位立法者得到 $1-z$,第三位立法者得到 0。得到 $1-z$ 的立法者如果被选中,就会推动这一提案通过。得到 0 的立法者如果被选中,就会提出修正案,留给自己 z ,给初始提案者

0, 给另一位立法者 $1-z$ 。从这一修正案中得到正的份额的立法者, 都会投票支持修正案。

此时, 要计算最优的 z , 我们必须考虑两个值。设 $v_i^1(z)$ 是在时期开始时, 使 i 得到 z 而另两位立法者分别得到 $1-z$ 和 0 的提案, 对立法者 i 的价值。设 $v_i^1(0)$ 为在时期开始时, 使 i 得到 0 而另两人分别得到 z 和 $1-z$ 的提案, 对立法者 i 的价值。由于我们只探讨对称均衡, 所以, 下标可以忽略。

我们首先计算 $v(z)$ 。以 $\frac{1}{2}$ 的概率, 这一提案得到通过, 提案者得到 z 。以 $\frac{1}{2}$ 的概率, 这一提案被修正, 于是在下个时期开始时提交的方案中, 初始提案者得到 0。因此:

$$v^1(z) = \frac{z}{2} + \frac{\delta v^1(0)}{2}$$

再看在时期开始时使 i 得到 0 的提案, 对 i 的价值。这一提案以 $\frac{1}{2}$ 的概率被通过, i 得到 0。 i 以 $\frac{1}{2}$ 的概率被随机选中, 对提案进行修正。于是, 下个时期开始时提交的方案中, 他得到 z 。因此:

$$v^1(0) = \frac{\delta v^1(z)}{2}$$

将这两个值结合在一起, 我们得到:

$$v^1(z) = \frac{z}{2} + \frac{\delta^2 v^1(z)}{4}$$

即

$$v^1(z) = \frac{2z}{4 - \delta^2}$$

最后还得保证得到 $1-z$ 的立法者愿意推动这一提案表决而不是对其进行修正。这要求 $1-z \geq \delta v^1(z)$, 即 $z \leq 1 - \delta v^1(z)$ 。提案者在这一约束条件下通过对 z 的选择, 最大化 $v^1(z)$ 。这一最大化问题的解为:

$$z = \frac{4 - \delta^2}{4 + 2\delta - \delta^2}$$

由此实现的持续值为:

$$v^1(z) = \frac{2}{4 + 2\delta - \delta^2}$$

要确定提案者所选择的策略, 我们只需比较 $v^1(z)$ 和 $v^2(z)$ 。计算结果表明, 在 $\delta > \delta^* \equiv \sqrt{3} - 1$ 时, $v^1(z) > v^2(z)$ 。就是说, 当参与者有耐性且重视未来时, 同时阻止两位立法者修正提案的代价很高。因此, 提案人只收买其中一位立法者, 而承受另一位立法者修正提案的风险。

这个开放规则模型有两个有意义的特点。首先,联盟的规模可能超过提案获得通过所需的最小规模。在 $\delta < \delta^*$ 时,就会出现这种情况:提案者让出了足够的资源以阻止对其提案的一切可能的修正。其次,可能存在均衡时滞。在 $\delta > \delta^*$ 且提案者留给其中一位立法者的份额为0时,就会出现这种情况。如果这位一无所获的立法者被随机选中,他就会提出能成功通过表决的修正案,而这实际上阻止了在第一个时期中达成一致。

我们把开放规则下的均衡配置与封闭规则下的均衡配置进行比较。有关文献对两种情况下提案者的所得做了很多研究。^①前面说过,在封闭规则下,当 $N = 3$ 时,提案者得到 $\frac{3-\delta}{3}$ 。这一份额始终大于 $v^2(z)$,并且在 $\delta > \delta^*$ 时大于 $v^1(z)$ 。就是说,开放规则降低了提案者的优势。超多数通过的投票规则还会削弱提案权。在 $k = N = 3$ 时,提案者得到的份额为 $\frac{3-2\delta}{3}$,始终低于 $v^1(z)$ 。就是说,在 $\delta > \delta^*$ 时,一致通过的投票规则把提案权降到了开放的多数通过的投票规则中没有高成本的滞后现象时的提案权水平之下。

10.5 不完全信息下的讨价还价

在此前讨论的所有讨价还价模型中,每个参与者都知道对手的分歧值或持续值,因此知道哪些方案会被接受,哪些方案会被拒绝。^②在许多政治问题中,这种假定显然不符合现实。立法机构不知道行政机构是否会签署某一法案;各个国家不知道和平条款是否会被接受,或者不知道对手是否想继续战争;如此等等。在本节中,我们介绍不完全信息下讨价还价模型的基本架构。后面两节则分别用行政机构和立法机构之间的讨价还价和国际关系中的危机谈判为例,阐述这一模型。

基础模型

在Nash的讨价还价问题中,设两个参与者就对 X 的分配在谈判。现在,在参与者 A 的分歧值上存在着不确定性。以概率 π ,这一分歧值为 $u_A = 0$;以概率 $1 - \pi$,为 $u_A = d > 0$ 。我们称 $u_A = 0$ 为“弱”类型,称 $u_A = d$ 为“强”类型。为使分析简单,假设参与者 B 向参与者 A 提出了一个或者接受或者不接受的分配方案 $(u_A, u_B) \in \Omega$ 。参与者 A 如果接受,则双方执行这一方案;参与者 A 如果拒绝,则双方得到的收益为 (u_A, u_B) 。

① 这不仅仅是个重要的公平问题。在分配高成本项目的收益的模型中,限制提案者的所得的制度削弱了通过无效率项目的激励(Baron, 1991; McCarty, 2000b; Primo, 2004)。

② 在开放规则下的Baron-Ferejohn博弈中,在哪个参与者会被选中而提出修正案的问题上,存在着不确定性;但是在某个参与者是偏好修正案还是偏好原始提案的问题上,没有不确定性。

显然,只有在分配方案能给他 $u_A \geq \underline{u}_A$ 时,参与者 A 才会接受这一方案。参与者 B 不知道 $\underline{u}_A$ 的值,所以他不知道需要给 A 多少效用才能使参与者 A 接受自己的方案。如果他给参与者 A 的效用低于参与者 A 的分歧值,参与者 A 就会拒绝接受,从而导致 $(\underline{u}_A, \underline{u}_B)$ 。设 $P(u_A | \underline{u}_A)$ 是参与者 A 在自己的分歧值为 $\underline{u}_A$ 时,接受 u_A 的概率。我们可以把参与者 B 的期望效用表示为其方案的函数:

$$EU_B(u_A) = [\pi P(u_A | 0) + (1 - \pi) P(u_A | d)] g(u_A) \\ + [1 - \pi P(u_A | 0) - (1 - \pi) P(u_A | d)] \underline{u}_B$$

其中, $g(u_A) = u_B(X - u_A^{-1}(u_A))$ 。参与者 A 的行为要具有序贯理性,要求在 $u_A \geq \underline{u}_A$ 时, $P(u_A | \underline{u}_A) = 1$; 否则,就等于 0。所以,我们可以把 B 的效用写为:

$$EU_B(u_A) = \begin{cases} \pi g(u_A) + (1 - \pi) \underline{u}_B, & \text{如果 } d > u_A \geq 0 \\ g(u_A), & \text{如果 } u_A \geq d \end{cases}$$

$g(u_A)$ 在 u_A 上递减,所以,唯一可能的解为 $u_A = d$ 或 $u_A = 0$ 。我们称 $u_A = 0$ 为进攻性分配方案,称 $u_A = d$ 为适应性分配方案。参与者 B 选择进攻性分配方案当且仅当:

$$\pi g(0) + (1 - \pi) \underline{u}_B > g(d)$$

即

$$\pi > \frac{g(d) - \underline{u}_B}{g(0) - \underline{u}_B}$$

这一模型简单,但是有很多合理的结论。首先,在以下条件得到满足时, B 更可能提出进攻性分配方案 $u_A = 0$:

- A 为弱类型的概率很高,
- B 的分歧值 $\underline{u}_B$ 很高,
- 进攻性方案和适应性方案之间的效用差 $g(0) - g(d)$ 很大。

由于假设了单边不完全信息和或者接受或者拒绝的分配方案,这一模型有很大局限性。我们不对这一抽象模型进行推广,而是探讨它在政治科学中的几个应用,在这些应用研究中,我们放松这些假设。

10.6 应用:包含否决行为的讨价还价

在第 7 章中,我们用 Romer-Rosenthal 的议程决定模型研究了总统否决权。这一不完全信息模型为分析否决权提供了相当好的工具,但是,它没有为研究否决行为提供基础;它认为,否决行为不会发生。我们现在用一个简单模型研究否决行为而不是否决权。在这一模型中,否决行为会发生。这一简单的不完全信息模型为建立包含信誉、学习和动态性的更复杂的否决谈判模型,提供了基础。^①

① 本节大量借用了 Cameron 和 McCarty(2004)。

为揭示否决行为会发生的事实,我们需要放松第7章中的基本模型的假设。在第7章的模型中,有很多限制性很强的假设,但是,对否决行为不会发生的结论,其中没有哪个假设真正起作用。但是,有一个例外,那就是C拥有关于P和O的偏好的完全信息。当立法机构面临不确定性时,否决行为可能发生,因为立法机构可能高估自己迫使总统或者关键投票人让步的能力。

放松完全信息假设,是近期很多否决谈判研究的起点(Matthews, 1989; McCarty, 1997; Cameron, 1999)。为展示这些模型的基本结构,我们假设不存在推翻总统否决令的可能性。所以,在总统否决之后,q依然为现行政策。为反映提案者C在接收者P的偏好上面临的不确定性,我们假设C认为P为两种偏好类型中的一种;这两种类型分别是:有理想点m的温和派,和有理想点e的极端派。我们依然假定所有参与者都有线性偏好:在政策为 $x \in \mathbf{R}$ 和理想点为 $i \in \{c, e, m\}$ 时,效用为 $-|x - i|$,其中, $e < m < c$ 。设P是极端类型的概率为 π 。

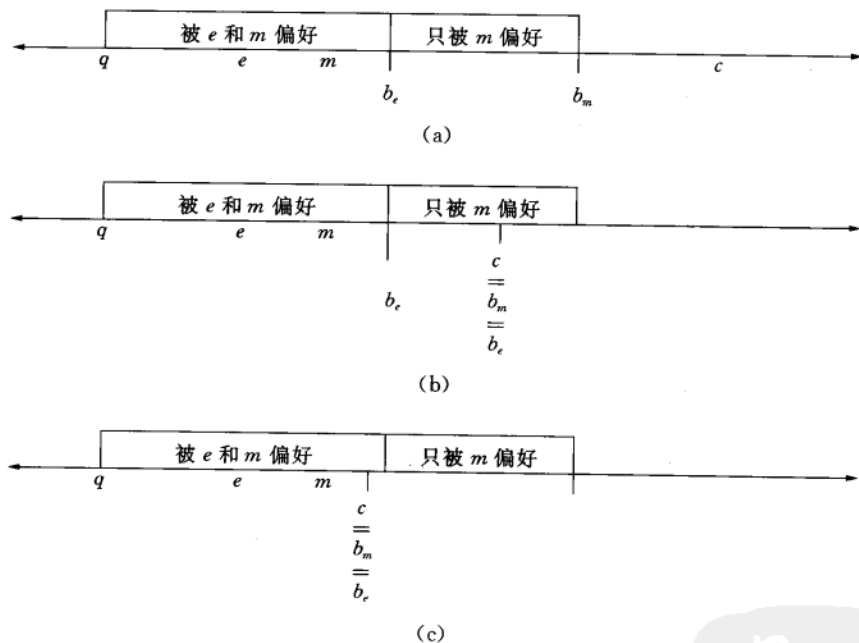


图 10.2 不完全信息下的否决谈判

由于偏好上的不确定性,C不再确切知道总统会接受或否决哪些法案。为说明这一点,我们看图10.2,其中 $q < e$ 。在这里,相对于现行政策,极端类型愿意接受的法案集只是温和类型愿意接受的法案集的子集。因此,C能迫使温和的接收者接受(从C的角度看)更好的法案;而在对方为极端类型时,她可能无法做到这一点。C面对的问题是,提出一份自己相对并不偏好但能被两种类型都接受的法案,例如法案 b_e ,还是更激进地提出一份她更偏好但只被温和类型接受的法案,例如法案 b_m 。C的选择显然决定于她对P的类型的推断。如果 π 很高(C认为P很可能

为极端类型), C 可能不会提出激进型法案。如果 π 很低 (C 认为 P 很可能是温和类型), C 就可能愿意赌上一把。如果她提出这样一份法案, 结果有可能很糟: P 是极端类型, 否决了这一法案。

我们来计算存在均衡否决行为的必要条件。假设图 10.2 中的偏好结构成立 (即 $q < e < m < c$)。设 $B_t(q)$ 是相对于现行政策, 被每一种类型 $t \in \{e, m\}$ 接受的法案集。和完全信息条件下的分析一样, 在 $t > q$ 时, 这一集合为 $[q, 2t - q]$; 否则, 为 $[2t - q, q]$ 。对任何的 q , 和总统 e 会接受的法案相比, 总统 m 接受的法案更多。由于 $e > q$, 所以,

$$B_e(q) = [q, 2e - q] \subset B_m(q) = [q, 2m - q]$$

就是说, 对 e 会接受的一切法案, m 都会接受; 但是, 这一结论的逆命题并不成立。所以, C 得有所取舍。 C 可以提出法案 $2e - q$, 两种类型的总统都会接受这一法案; 或者提出法案 $2m - q$, 总统 e 会否决这一法案。给定 C 的推断, 后一法案以 π 的概率被否决。

情况 1: $c > 2m - q$ 。给定 C 的线性偏好, C 从 $b = 2e - q$ 中得到的效用为 $2e - q - c$; 从 $b = 2m - q$ 中得到的效用为 $\pi q + (1 - \pi)(2m - q) - c$ 。如果 $\pi \leq \frac{m - e}{m - q}$, C 偏好 $b = 2m - q$, 尽管它将以概率 π 被否决。

情况 2: $2e - q < c < 2m - q$ 。 C 从 $b = 2e - q$ 中得到的效用为 $2e - q - c$ 。现在, m 接受法案 $b = c$ 。于是, C 如果提出自己的理想点, 得到期望效用 $\pi(q - c)$ 。所以, 在 $\pi \leq \frac{c + q - 2e}{c - q}$ 时, C 会提出法案 $b = c$ 。在这里, π 的临界值小于情况 1 中的临界值, 这使得在这一偏好结构下, 法案不大可能被否决。

情况 3: $c < 2e - q$ 。现在, 两种类型的总统都接受法案 $b = c$ 。所以, 对一切的 π 值, C 都提出自己的理想点, 法案不会被否决。

这一简单模型的基本结论是, 当 C 的偏好接近 m 和 e 时, 法案不大可能被否决。Cameron(1999)的经验分析发现, 在美国国会和总统由同一个党派控制时, 法案被否决的可能性下降: 他认为这一事实佐证了上述结论。

10.6.1 包含信誉、学习和动态性的模型

这个不完全信息模型有一个特点: 在 P 是温和型时, 如果提案者 C 误认为 P 是极端型, 温和型总统就能得到更好的结果。这就存在一种可能性: P 可能会操控 C 关于类型 (即信誉) 的推断。在本节将分析的三个模型中, 参与者试图操控 P 的信誉。这三个模型都是信号揭示模型, 因为拥有更多信息的一方通过行动传递关于 P 的类型的信息。在前两个模型中, 即在否决威胁模型和序贯否决谈判模型中, 拥有更多信息的一方是总统。在第三个模型中, 即在“谴责博弈”否决模型中, C 和 P 都通过行动向缺乏信息的投票者传递信息。

10.6.1.1 否决威胁

从有戏剧色彩的各种“暗示”，到美国政府管理和预算办公室定期发布的“政府政策咨文”等，否决威胁构成了美国立法政治的重要特点。上述各种模型没有为理解这一现象提供帮助。Matthews(1989)提出了一个很有影响的否决威胁模型。在他的模型中，总统利用“空谈”揭示关于自己的偏好和否决意图的信息。

为更好阐述这一模型，我们把总统类型从两种增加到四种：除 m 和 e 外，我们用 r 表示顽固型， a 表示适应型。每种类型的总统和 C 都有线性偏好，并且， $r < q < e < m < c < a$ ，见图 10.3。总统 r 被称为顽固型，因为他否决 C 认为优于现行政策的一切法案。总统 a 是适应型的，因为相对于现行政策，他偏好 c 。各种类型的发生概率分别为 π_r 、 π_e 、 π_m 和 π_a 。在这一博弈中，总统首先“讲话”（向立法机构发出无成本信号）。他所传递的信息没有实际意义，而仅是源自其均衡策略的背景性信息。在总统讲话后， C 更新自己关于总统偏好的推断，然后提出被总统接受或否决的提案。

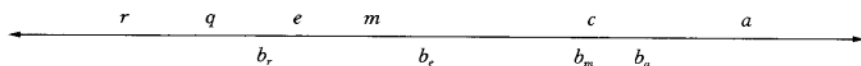


图 10.3 “空谈”否决谈判模型中的最优反应

首先假设在均衡中，总统讲话没有揭示任何信息，因为在可能的讲话内容构成的集合上，各种类型的总统使用相同的混合策略。^①在这一废话均衡中， C 从 $b_r = q$ ， $b_e = 2e - q$ ， $b_m = 2m - c$ 和 $b_a = c$ 中选择一份提案来最大化自己的效用。其中的任何一项选择，更低类型的总统会否决。例如，如果 C 选择 b_m ，类型为 e 和 r 的总统会否决，所以，这一法案被否决的概率是 $\pi_e + \pi_r$ 。这里不对每项法案的成立条件给出相应公式，而是在图 10.4 中用图形说明。给定 π_a 和 π_r 值，对 (π_m, π_e) 的各种取值，第一幅图给出了废话均衡中相对应的各种提案。 C 不会提出法案 $b_r = q$ ，因为 C 从被否决的提案中能得到和 $b_r = q$ 相同的效用。从总统角度看，这一废话均衡有些差。如果总统为类型 a ， C 如果提出政策 b_m 和 b_e ，总统的效用没有达到最好状态。如果总统为类型 m ，如果 C 提出法案 c （总统 m 会否决这一提案）而不是总统所偏好的 b_e ，总统的效用也没有达到最好状态。总统 r 和总统 e 则从所有提案中得到当前效用水平，所以，他们不受影响。 C 也受困于信息匮乏，因为这可能迫使 C 或者做出不必要的让步或者就得面对着法案被否决的风险。

鉴于废话均衡中的结果不理想，我们自然要问，是否还存在其他均衡，有更多的信息被传递出来。Matthews 指出，关于总统偏好的一些信息，但不是全部信息，能够在总统的讲话中揭示出来。我们首先探讨，不同类型总统发表不同讲话的分

① 在所有的讲话内容上的这种随机选择意味着没有偏离均衡的历史。这里不存在所有类型的总统给出相同讲话内容的至圣废话均衡。请你证明这一结论。

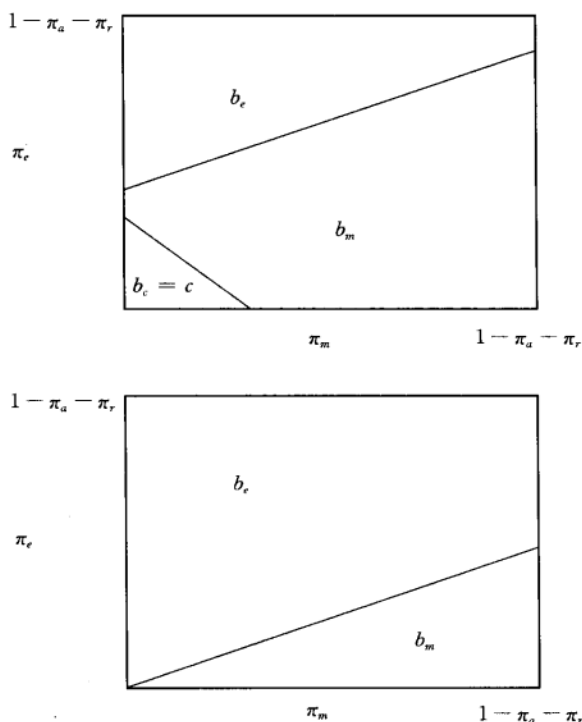


图 10.4 上图:“废话均衡”中的提案 下图:在妥协信号后,“两信号均衡”中的提案

离均衡,为什么不可能存在。如果 C 能够从讲话中推断出总统的类型, C 就会最优地对 r 提出法案 b_r , 对 e 提出法案 b_e , 如此等等。但是, 由于在 b_e 和 b_m 之间, m 偏好 b_e , 所以 m 会选择背叛, 即发表 e 的讲话。所以, 分离均衡不可能存在。Matthews 指出, 在能揭示出最多的信息的均衡中, 类型 a 发表“适应性”讲话, 其他所有类型发表相同的“威胁性”讲话。在听到适应性讲话后, C 正确地推断出总统会接受自己的理想点, 于是提出法案 c 。类型为 a 的总统之所以发表适应性讲话, 是因为他偏好 c 而不是 b_m 或 b_e 。在听到威胁性讲话后, C 知道总统类型不是 a , 于是相应地更新自己的推断, 并在 b_m 和 b_e 之间做出选择。给定 π_a 和 π_r 的值, 图 10.4 的第二幅图把最优法案表示为 π_m 和 π_e 的函数。这里有两点需要指出。首先, C 更可能提出法案 b_e , 因为总统类型不是 a , 所以 b_m 遭到否决的概率更高。所以, 和听到适应性讲话相比, 在听到威胁性讲话后, C 对总统的偏好做出了更大让步。其次, 这种能最大程度地揭示信息的均衡未必存在。假设类型 a 在 b_m 和 c 之间偏好 b_m , 在 c 和 b_e 之间偏好 c 。只有在 C 对威胁性讲话的最优反应是 b_e 时, 这种均衡才存在; 在其他情况中, a 会选择威胁性讲话。同样地, 如果在 b_e 和 c 之间, a 偏好 b_e , 就不存在这种均衡。

对某些偏好结构, 否决威胁不起作用。设想在图 10.3 中, m 向右移动使得其在 c 和 q 之间, m 偏好 c (就是说, 他变为适应型), 但是在 b_e 和 c 之间, 他依然偏好

b_e 。故在能最大程度地揭示信息的均衡中, m 发表威胁性讲话,但是,这种威胁不起任何作用,因为他会接受 C 的理想点。

这种能揭示出信息的均衡对 C 有利(否则, C 可以关掉电视机,不在乎总统的讲话)。但是,某些类型的总统的效用可能下降。调整 a 的位置,使得在 b_m 和 c 之间,他偏好 b_m ,在 c 和 b_e 之间,他偏好 c 。再假设废话均衡产生结果 b_m ;在能最大程度地揭示信息的均衡中,威胁性讲话产生结果 b_e 。 a 显然偏好废话均衡中的结果。

10.6.1.2 不完全信息下的序贯否决谈判

提案人常常会提出多个法案,在这一过程中,提案人能了解到接收者的类型。例如,如果接收者曾经否决过路线强硬的提案,提案人就可以认为接收者确实是强硬类型的。于是,提案人的下一个提案可能更适应接收者的类型。在许多形式的讨价还价中,这种动态特点十分常见;在否决谈判中也能看到这种现象。但是,情况也可能正好相反:接收者常常会为塑造一种形象而否决前期提案,借此在博弈后期得到更好的提案。但是,认识到这一点,提案者为什么还应该在提案上做出妥协呢?序贯否决谈判模型探讨此类与学习和可信性等有关的问题。

我们看一个简单例子,它揭示了这类模型的许多基本思想。设 $q = 0$, $e = 0.25$, $m = 0.6$, $c = 1$ 。显然,在(没有否决威胁的)一次性博弈中, $b_e = 0.5$, $b_m = c$ 。利用一次性不完全信息模型的这些结论,我们能够证明,如果 $\pi < \frac{1}{2}$, C 提出法案 $b_m = c$; 否则,提出法案 $b_e = 0.5$ 。假设这一博弈不是一次性的,即在第一份法案被否决后, C 可以提出第二份法案。假设双方的谈判以 ρ 的概率破裂,如果谈判没有破裂,允许 C 提出第二份法案。谈判破裂的可能性反映了立法过程和其他政治过程中内在的不确定性。那么,是否存在如下的讨价还价均衡: C 首先提出立场强硬的提案,在此法案被否决后但双方间的讨价还价没有破裂时,再提出一份更适应对方类型的提案吗?

在这种讨价还价均衡中,温和型总统肯定在第一轮中接受立场强硬的提案(如果两种类型的总统都否决这一提案,则 C 应该提出适应性提案,避免谈判破裂而使自己面对比提案糟糕的现状)。因此,下面的激励一致性约束必须得到满足:

$$m - c \geq (1 - \rho)(m - 2e + q) + \rho(q - m)$$

即

$$m - \rho \geq \frac{c - q - 2(m - e)}{2(e - q)}$$

这一激励一致性约束指出,考虑到谈判破裂的可能性,对温和型总统来说,在第一轮中就接受强硬的提案,比否决此提案和等待下一个适应性提案,更加有利。在这个例子中,谈判破裂概率的临界值是 0.6。在提案遭到否决后, C 推断总统为极端的概率是 $\mu(e)$ 。这里要指出的是:在讨价还价均衡中, $\mu(e) \geq \frac{1}{2}$; 否则,在遭遇

否决后, C 会在最后一个阶段再次提出强硬提案(前面证明过这一结论)。但是, 如果谈判破裂的概率超过 0.6, 温和型总统就会接受初始提案, 于是, 根据贝叶斯法则, $\mu(e) = 1$; 在否决后, C 在第二轮中提出适应性提案。

我们还需分析另一个激励一致性约束条件。在国会看来, 首先提出强硬提案, 然后(在前一提案被否决但双方间的谈判没有破裂时)提出适应性提案, 相对于在博弈开始时就提出肯定被接受的适应性提案, 必须更加有利可图。这就要求:

$$\pi[(1-\rho)(2e-q-c) + \rho(q-c)] \geq \rho(2e-q-c)$$

即

$$\pi \leq \frac{c-2e-q}{c-2e(1-\rho)+q(1-2\rho)}$$

在这一例子中, 这一条件为 $\pi \leq \frac{1}{1+\rho}$ 。现在, 在两阶段且有两种类型的序贯否决谈判模型中, 如果理想点如上所述, 我们就能找出它的讨价还价均衡。在 $0.6 \leq \rho \leq 1$ 且 $\pi \leq \frac{1}{1+\rho}$ 时, C 提出的提案为 $b_1 = b_m = c$ 和 $b_2 = b_e = 0.5$ 。类型为 m 的总统接受 b_e 和 b_m , 类型为 e 的总统接受提案 b_e , 否决提案 b_m 。在提案被否决后, C 推断总统为极端型的概率是 $\mu(e) = 1$ 。

10.6.1.3 多个法案上的讨价还价

上一节指出, 不完全信息能够影响单一法案上的讨价还价的动态特征。McCarty(1997)探讨了信息激励和信誉激励如何影响多个法案上的序贯谈判。在他的不完全信息否决谈判模型中, P 和 C 就一系列政策讨价还价, 现状点则为 q_1 和 q_2 。在两个时期中的每个时期里, C 提出法案 b_i , 总统决定接受还是否决这一法案。于是, 每个政策上的讨价还价被看作是个一次性博弈。如果 P 否决 b_i , 结果就为现行政策 q_i 。总统的理想点被假定保持不变, 于是, 在 C 就政策 2 提出提案前, 政策 1 上的结果可能为 C 提供信息。在最后一个时期中, 博弈和前面所描述的一次性不完全信息博弈完全一样, 因此, 如果偏好为图 10.2(a) 和图 10.2(b) 中的结构, 则类型为 m 的总统, 如果能让 C 认为自己为极端型, 就能使自己在第二项政策上得到更好的结果。就是说, 给定图中的偏好结构, m 为在政策 2 上得到更好的结果, 可能在第一个时期中否决法案, 把自己塑造为极端型总统。这就意味着他得否决相对于 q_1 自己更偏好的法案; 类型为 e 的总统不会这么做。所以, 信誉激励提高了政策 1 被否决的概率。 C 认识到这种动机的存在, 可以在第一个政策上做出足够大的让步, 从而阻止类型为 m 的总统为塑造形象而行使否决权。McCarty 的模型预测了“蜜月”模式的存在——在总统任期之初, 国会提出适应性政策; 随着总统任期趋于结束, 在信誉动机逐步消失时, 提出较少让步的政策。但是, 他指出, 信誉激励的存在决定于偏好结构, 如图 10.2(a) 和 (b) 中的偏好结构, 所以, 当 P 和 C 之间的预期分歧很小时, 如在国会和总统都由一党控制时, 蜜月效应不大可能出现。

10.6.2 谴责博弈

Groseclose 和 McCarty(2000)指出,否决行为更多地是选举政治而不是立法过程中的不确定性的产物。在他们的模型中,提案人利用所掌握的提案权,向外界揭示总统政策与选民利益相悖。在这种“谴责”博弈模型中,当提案人从揭示总统有极端偏好中得到的收益,大于其从推行新政策中得到的收益时,否决行为就会发生。就是说,重要的是选民在对总统的认识上的不确定性,而立法者面对的不确定性并不重要。

我们看这一模型的一个简化版本。这一模型中有一位新的参与者 V , 即选民。他有线性偏好,理想点为 v 。我们依然使用上一节中的有关前提。 V 以概率 π 认定总统为类型 e , 以概率 $1-\pi$ 认定其为类型 m 。我们只探讨 $e < m < v$ 时的情况。假设选民根据总统的理想点与自己的理想点之间的期望距离评价总统表现。就是说,选民对总统的评价为:

$$w(e, m, \pi; v) = -\pi |v - e| - (1 - \pi) |v - m| = \pi e + (1 - \pi)m - v$$

这一模型的一个重要特点是, P 和 C 都关心 V 从总统的立场中得到多少期望效用。最有意义的情况出现在他们之间存在着利益冲突时——当选民认为总统是温和型时,总统得到更大的效用;当选民认为总统是极端型时,国会得到更大的效用。在国会和总统职位由不同政党或团体控制时,特别当这些政党高度极化且中间选民为温和型时,这种情况就可能出现。在这种情况下, C 和 P 就得在从政策中得到的收益和从政治姿态中得到的收益之间,进行权衡取舍。具体地说,总统偏好能导致公众降低 π 的那些行动;立法机构偏好能导致公众增加 π 的那些行动。我们用 λ_c 和 λ_p 分别表示 C 和 P 赋予政策的权重,借此反映他们对这两种收益不同的重视程度。于是, C 和 P 的效用函数分别为:

$$-\lambda_c |x - c| + (1 - \lambda_c)(\pi e + (1 - \pi)m - v)$$

和

$$-\lambda_p |x - p| - (1 - \lambda_p)(\pi e + (1 - \pi)m - v)$$

这里的一个重要假设是, V 并不了解 P 的偏好, C 对此却有充分信息。所以, C 能够通过对法案的选择,把自己拥有的关于 π 的信息传递给选民。总统接受或否决某些法案的决策同样可能向选民揭示关于其偏好的信息。

一个特别有意思的均衡是,在 P 为温和型时, C 提出可被 P 接受的法案;在总统为极端型时,提出会被 P 否决的法案。当且仅当以下两个条件成立时,这一均衡存在:

$$\frac{\lambda_p - \lambda_c}{\lambda_p \lambda_c} (1 - \pi)(m - e) \geq 2(e - q) \quad (10.1)$$

$$2 \geq \frac{\lambda_p - \lambda_c}{\lambda_p \lambda_c} \pi \quad (10.2)$$

这两个条件给出了关于否决行为发生与否的一些结论。^①首先,如果 $m = e$ 或者 $\pi = 1$, 条件(10.1)不可能得到满足。所以,选民在总统偏好上的不确定性是关键。如果没有这种不确定性,法案被否决的决策,对 C 而言,失去了信号价值,于是,在这种情况下, C 愿意向两种类型的总统提交能被接受的法案。其次,在 π 比较小时,这两个条件容易得到满足。由于对总统的事前评价在 π (即总统为极端型的概率)上递减,所以这一模型预测,当公众认为总统是温和型时(即认为总统在政策立场上与自己接近),否决行为会更经常地发生。直观地说,当国会掌握关于总统的政策偏好的负面信息,而这种信息又与选民的推断相矛盾时,这种谴责博弈对国会更有用。

下面三个结论所依托的基础是, C 和 P 愿意用政策收益交换政治收益。图 10.5 说明了政策权重 λ_p 和 λ_c 如何影响这两个条件,上方曲线以下的区域为满足条件(10.1)的 λ_p 和 λ_c 的组合。下方曲线以上的区域为满足条件(10.2)的 λ_p 和 λ_c 的组合。谴责博弈均衡位于这两个区域的交集中。首先,只有在 $\lambda_p > \lambda_c$ 时,就是说总统赋予政策结果的权重比国会赋予的权重更大时,条件(10.1)才得到满足。否则, C 更愿意通过双方都偏好的法案,借此获得政策收益,而不是通过会被总统否决的法案而获得的选举优势。条件(10.2)对政策权重上的分歧施加了上界。如果 λ_p 比 λ_c 大得多, C 就无法用自己的提案发出可信赖的信号的能力。最后一个结论来自这一事实:在谴责博弈模型中,只有极端型总统才会否決法案。由于只有类型为 e 的总统否決法案,所以,否決法案的决策,导致选民支持率下降。

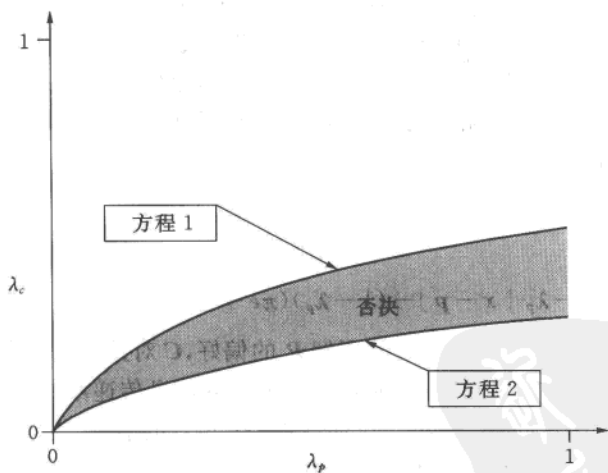


图 10.5 在谴责博弈模型中,均衡否决行为的存在条件

① 当 $c > 2m - q + \frac{(m-e)(1-\lambda_p)}{\lambda_p}$ 时,这两个条件构成必要条件。 c 的不同位置产生略有不同的但在性质上相同的条件。

10.7 应用:危机谈判

讨价还价理论的一个局限性是,它的解一般高度依赖谈判规则,因此易于受扩展形式上的变化影响。在否决谈判模型中,这不是个大问题,因为它的规则常根植于宪法条款和明确的立法程序中。但是,在主权国家之间的危机谈判中,让讨价还价博弈的解高度依赖博弈的某些扩展形式,显然并不十分合理;此时适用的规则,由于保密等因素,一般是不正式的、非条款化的和无法被观察到的。

认识到这一问题,Banks(1990)探讨了许多危机谈判博弈的均衡必须共同拥有的均衡特征。我们看下面这种危机谈判博弈。国家1和国家2就一个单位的领土在谈判。设 x 是国家1得到的份额。遵循Banks的分析,我们假定这两个国家风险中性,于是,国家1从谈判中得到的收益是 x ,国家2的收益是 $1-x$ 。如果双方无法就这块领土的分割达成一致,就会导致战争。^①战争带给国家1的期望效用是 u ,带给国家2的期望效用是 v 。这些期望效用包含了对赢得战争的概率、从战胜和战败中得到的收益和战胜国能够得到的领土份额等因素的预期。

Banks假定,相对于国家2,国家1在 (u, v) 的值上拥有信息优势。Banks使用常用的分析方式处理不对称信息的处理。他假定国家1的类型为某个 $t \in T$,国家1从战争中得到的期望收益 $u(t)$ 在其类型上递增。国家1在谈判前就了解 t 的大小,而国家2只有先验推断 $f(t)$, $f(t)$ 是公共知识。

给定这一分析框架,标准博弈论模型赋予每个国家一个决策集和一个(概率)结果函数,这个函数明确了发生战争的概率,并把最终解决方案的分布定义为决策的函数。从这一模型中,我们能够得到均衡策略 $(\sigma_1(t), \sigma_2)$,这一策略产生均衡战争概率 $p(t)$ 和期望解决方案 $x(t)$ 。在这种均衡中,国家1的收益是:

$$U(t; x, p) = p(t)u(t) + (1 - p(t))x(t)$$

显然,不是所有的 $p(t)$ 和 $x(t)$ 都能从贝叶斯均衡中得到。具体地说,贝叶斯均衡需要两个条件。第一是激励一致性条件。在结果 $(p(t'), x(t'))$ 和 $(p(t), x(t))$ 之间,类型 t 不能偏好前者。否则,它就会偏离自己的均衡策略 $\sigma_1(t)$ 。就是说,对每个 t 和 t' ,激励一致性要求:

$$p(t)u(t) + (1 - p(t))x(t) \geq p(t')u(t) + (1 - p(t'))x(t') \quad (10.3)$$

$$p(t')u(t') + (1 - p(t'))x(t') \geq p(t)u(t') + (1 - p(t))x(t) \quad (10.4)$$

贝叶斯均衡所施加的第二个条件是个人理性。^②除非 $p(t) = 1$,否则 $u(t)$ 不能大于 $x(t)$ 。否则, t 就会撕毁协议和以1的概率发动战争。个人理性约束为:

① 在可能达成的协议中,包括维持现状。

② (在下一章,)你会发现 Banks 的分析方法与机制设计十分相似。

$$p(t)u(t) + (1 - p(t))x(t) \geq u(t) \quad (10.5)$$

激励一致性和个人理性显然只是给出了最基本的约束条件,但是,它们对谈判结果产生了相当多的影响。贝叶斯均衡的最重要的特征与 p 、 x 和 U 在 t 上的单调性有关。

引理 10.1 如果 p 和 x 满足激励一致性和个人理性,则 $p(t)$ 在 T 上弱递增。

证明: 设 $t' > t$ 。(10.4)式左边部分减去(10.3)式右边部分,(10.4)式右边部分减去(10.3)式左边部分,得到:

$$p(t')[u(t') - u(t)] \geq p(t)[u(t') - u(t)]$$

由于 $u(t)$ 严格递增,所以, $u(t') - u(t) > 0$, 于是, $p(t') \geq p(t)$ 。□

这一引理指出,在任何贝叶斯均衡中,发生战争的概率不可能随国家 1 的战争期望效用的上升而下降。这不仅是策略模型的特点,而且与关于战争的许多决策理论模型的假定相一致。

为探讨下一个结论,我们定义对给定的 p 和 x ,能以正概率通过谈判解决双方间冲突的类型集: $T_b = \{t \in T: p(t) < 1\}$ 。在这里,对任何 $t \in T_b$, 个人理性约束条件要求 $x(t) \geq u(t)$ 。

引理 10.2 如果 p 和 x 满足激励一致性和个人理性,则 $x(t)$ 在 T_b 上弱递增。

证明: 设 $t, t' \in T_b$ 且 $t' > t$ 。根据引理 10.1, $1 > p(t') \geq p(t)$ 。由于对所有的 $t \in T_b$, $x(t) \geq u(t)$, 所以,

$$p(t)u(t') + (1 - p(t))x(t') \geq p(t')u(t') + (1 - p(t'))x(t') \quad (10.6)$$

将这一结论与(10.4)式结合,得到:

$$p(t)u(t') + (1 - p(t))x(t') \geq p(t)u(t') + (1 - p(t))x(t)$$

不等号两边同时除以 $1 - p(t)$ —— 由于 $t \in T_b$, 所以, $1 - p(t)$ 为正 —— 于是得到 $x(t') \geq x(t)$ 。□

国家 1, 在其战争效用上升时, 从谈判结果中得到的效用至少不能变差。引理 10.1 和引理 10.2 指出, 在任何贝叶斯均衡中, 更高的类型得到更好的谈判结果, 但是发生战争的风险更大。^①但是, 激励一致性要求这些替代关系有利高类型(否则, 他们就会把自己伪装为低类型)。设 $T_w = \{t \in T: p(t) > 0\}$, 就是说, T_w 是以正概率进行战争的所有的类型构成的集合。

引理 10.3 如果 p 和 x 是激励一致的和个人理性的, 则 $U(t; x, p)$ 在 T 上弱递增, 在 T_w 上严格递增。

证明: 设 $t' > t$ 且 $t', t \in T_w$ 。(使用反证法)假设 $U(t) \geq U(t')$, 则有:

$$p(t)u(t) + (1 - p(t))x(t) \geq p(t')u(t') + (1 - p(t'))x(t')$$

① Banks 还指出, 如果 x 和 p 满足激励一致性和个人理性, 则 $x(t') > x(t)$ 当且仅当对 $t, t' \in T_b$ 且 $t > t'$, $p(t') > p(t)$ 。

由于 $u(t') > u(t)$, 因此:

$$p(t)u(t') + (1-p(t))x(t) \geq p(t')u(t') + (1-p(t'))x(t') \quad (10.7)$$

由于 $t', t \in T_w$, $p(t)$ 和 $p(t')$ 大于 0, 因此, (10.7) 式与 (10.4) 式矛盾。所以, $U(t') > U(t)$ 。在集合 $T \setminus T_w$ 上, 这一严格不等式不成立。如果 $t, t' \in T \setminus T_w$, 则 $p(t) = p(t') = 0$, 根据 (10.3) 式和 (10.4) 式, $x(t) = x(t')$, $U(t') = U(t)$ 。□

所以说, 战争的期望效用上升, 不会使国家 1 福利下降。事实上, 对以非零概率进行战争的国家, 更高的战争收益导致更高的均衡收益。^①

激励一致性分析方法虽然能告诉我们危机谈判模型共有的一些特征, 但是, 依然有很多问题是它无法解决的。例如, 我们不知道国家 2 对国家 1 的战争效用的认识, 即先验推断 $f(t)$, 如何影响发生战争的概率或谈判解决冲突的可能性。对这些问题, 我们求助于更具体的危机谈判模型。

10.7.1 危机谈判模型

Fearon(1995)探讨了几个危机谈判模型, 它们都在 Banks 所分析的博弈类型中。Fearon 给出了 (u, v) 具体的形式。假定每个国家的战争成本为 $c_i > 0$ 。国家 1 以概率 $\pi \in (0, 1)$ 赢得任何战争。战胜方得到它最偏好的解决方案(对国家 1, 是 $x = 1$; 对国家 2, 是 $x = 0$)。于是, $u = \pi - c_1$, $v = 1 - \pi - c_2$ 。设 x_0 是争议地区的当前分配格局。

在这一框架下, 我们可以定义为避免战争可为各个国家接受的协议集。对国家 1, 我们要求 $x = \pi - c_1$; 对国家 2, 我们要求 $1 - x > 1 - \pi - c_2$ 。因此, 任何配置 $x \in [\pi - c_1, \pi + c_2]$ 都能阻止战争。由于战争成本为正, 所以, 和平协议集为非空集。我们假定, 在完全信息下, 存在某个协议能被双方接受, 战争得以避免。^②

在 Fearon 的模型中, 在国家 1 的成本上存在着不完全信息。他假定 c_1 有连续分布, 但是, 我们假定 c_1 只取两个值: $\underline{c}$ 和 $\bar{c}$ 且 $\bar{c} > \underline{c}$ 。作为共同知识的先验推断是: 以概率 λ , $c = \underline{c}$ 。在 Fearon 探讨的第一个模型中, 国家 2 向国家 1 提出一份或者接受或者拒绝的分配方案。如果国家 1 拒绝这一方案, 战争随即爆发。

在这一模型中, 如果国家 2 提出 $x \geq \pi - \underline{c}$, 两种类型的国家 1 都会接受这一方案, 战争得以避免。国家 2 显然没有动机让出高于 $\pi - \underline{c}$ 的份额给国家 1, 所以, 设 $\underline{x} = \pi - \underline{c}$ 。如果国家 2 提议 $x \in (\pi - \bar{c}, \pi - \underline{c}]$, 则只有低类型的国家 1 接受, 战争以 λ 的概率发生。在这些提议中, 国家 2 偏好 $\bar{x} = \pi - \bar{c}$ 。如果国家 2 提议 $x < \pi - \bar{c}$, 两种类型的国家 1 都会拒绝, 战争肯定爆发。这种情况下, 国家 2 得到的收益是 $v = 1 - \pi - c_2$ 。所以, 国家 2 实际上要在 $1 - \underline{x}$ 、 $\lambda v + (1 - \lambda)(1 - \bar{x})$ 和 v 三个效用

① Banks 同时证明了 $U(t; x, p)$ 在 t 上连续。我们对此不展开讨论, 建议读者阅读 Banks 的原始论文。

② 这实际上是在假定, 争议地区可以被无限分割, 从而使这一区间中的任何协议都可行。在一些冲突中, 这一点未必成立。当所争议的地区涉及宗教或理想认同等因素时, 这类问题就可能出现。

中做出选择。我们可以把第三个选项排除掉。因为区间 $[\pi - \underline{c}, \pi + c_2]$ 是非空集, 所以, $1 - \underline{x} > v$ 。就是说, 国家 2 永远不会破坏谈判而以 1 的概率发动战争。我们来确定国家 2 在剩下两个选项上的偏好。显然, 只要:

$$\lambda \leq \frac{\bar{c} - c}{c_2 + \bar{c}}$$

就有 $\lambda v + (1 - \lambda)(1 - \bar{x}) \geq 1 - \underline{x}$ 。当国家 1 为低成本类型的概率足够小时, 国家 2 就会采取强硬的谈判姿态——在国家 1 为低成本类型时, 国家 2 的这种姿态会引发战争。这里要注意的是, λ 的临界值在国家 2 的成本上递减; 在它的军力有限时, 它不愿意冒战争风险。

我们能够证明, Banks 的结论在这一模型成立。当 $\lambda > \frac{\bar{c} - c}{c_2 + \bar{c}}$ 时, 两种类型的国家 1 得到相同的份额, 发生战争的概率为 0; 当 $\lambda \leq \frac{\bar{c} - c}{c_2 + \bar{c}}$ 时, 两种类型的国家 1 得到相同的份额, 但是, 类型为 $\bar{c}$ 的国家 1 会以 1 的概率发起战争。

Fearon 指出, 对战争的这种信息解释, 我们有理由持怀疑态度。双方的信息沟通应该能解决这种信息不对称性和阻止战争。但是, 如果两个国家有截然相反的偏好, 我们能够证明, “cheap talk” 并不影响双方的谈判或发生战争的概率。设国家 1 发出信号 H 或 L , 借此揭示其成本为 $\bar{c}$ 或 c 。^① 观察到这些信号, 国家 2 更新自己对国家 1 的成本的推断, 新的推断用 λ^* 表示。国家 2 会根据这些更新后的推断, 使用和此前相同的截点规则。我们首先探讨是否存在分离均衡, 在其中, 类型 $\bar{c}$ 发出信号 H , 类型 c 发出信号 L 。在这种情况下, $\lambda^*(H) = 0$, $\lambda^*(L) = 1$, $x(H) = \bar{x}$, $x(L) = \underline{x}$ 。发生战争的概率是 0。分离均衡要求激励一致性条件得到满足: $x(H) \geq x(L)$, $x(L) \geq x(H)$ 。由于 $\underline{x} > \bar{x}$, 这两个不等式显然并不成立。我们请读者证明, 在这一博弈中, 不存在部分揭示信息的半混同均衡。

10.7.2 危机持续模型*

本节以 Fearon(1994)为基础。在这篇文章中, Fearon 利用消耗战模型, 探讨在国际冲突中落败国所承受的“观众成本”如何影响危机的动态性。

国家 1 和国家 2 为瓜分价值为 $v > 0$ 的利益而陷入冲突。这一博弈在连续时间上进行, 开始时点为 $t = 0$ 。在每个时点上, 每个国家都有三种策略可供选择: 进攻、放弃或持续。这个博弈一直进行下去, 直至其中一个或两个国家选择放弃或选择进攻。如果两个国家都选择持续, 这个博弈就继续进行。如果一个国家在对方选择放弃之前进攻, 双方从战争中得到的期望收益为 $w_i < 0$ 。Fearon 把这些收益解释为冲突的终端解决方案; w_i 更高的国家, 更愿意用军事冲突解决双方冲突。

① 是否只使用这两种信号或者是否赋予这两种信号实际含义等, 并不重要。

Fearon 所分析的是,对在冲突中落败的国家领导人施加的惩罚,如何影响危机行为。他假定在时点 t , 如果国家 i 先于国家 j 选择放弃, 它得承受观众成本 $a_i(t)$, 这一成本在 t 上严格递增。成本对 t 的依赖反映了这样一种认识: 相对于暂时冲突, 退出持续冲突的成本更高。

这一博弈的纯策略, 为在时点 t' 开始的子博弈规定了行动准则, 即规定了在每个有限时点 $t \geq t'$ 上, 是进攻还是放弃。^① 我们把这些策略表示为 $\{t, \text{进攻}\}$ 或 $\{t, \text{放弃}\}$, 分别表示“升级直至时点 t , 然后进攻”和“升级直至时点 t , 然后放弃”。

在探讨双方都不了解彼此的解决方案的更一般模型之前, 我们探讨完全信息模型。对每个国家, 我们算出它严格偏好进攻而不是放弃的时点 t 。显然, 在 $w_i \geq -a_i(t)$ 时, 这种情况就会出现。设:

$$\bar{t}_i \geq -a_i^{-1}(w_i)$$

由于国家 1 的解决方案或其观众成本高于国家 2, 所以, $\bar{t}_1 < \bar{t}_2$ 。就是说, 在 $\bar{t}_1$, 国家 2 愿意放弃以免被进攻; 因此, 它选择放弃。设 $Q_i(t; t')$ 是国家 i 在时点 t' 之前没有放弃的前提下, 在时点 t 之前选择放弃的概率; $Q_i(t)$ 是在时点 t 之前就选择放弃的无条件概率。

因此, 在每个子博弈 $0 \leq t' < \bar{t}_1$ 中, 国家 2 如果选择立刻放弃, 得到的收益为 $-a_2(t')$; 如果选择 $\{t, \text{放弃}\}$, 得到的收益为 $Q_1(t; t')v - (1 - Q_1(t; t'))a_2(t)$ 。我们看国家 1 的策略。对子博弈 $0 \leq t' < \bar{t}_1$, 策略 $\{\bar{t}_1, \text{进攻}\}$ 带给它的收益为 v , 策略 $\{t, \text{放弃}\}$ 带给它的收益为 $Q_2(t; t')v - (1 - Q_2(t; t'))a_2(t) < v$ 。因此, 对所有的 $t < \bar{t}_1$, $Q_1(t; t') = 0$ 。现在可以看出, 国家 2 从 $\{t, \text{放弃}\}$ 中得到的收益为 $-a_2(t) < -a_2(t')$ 。因此, 在每个子博弈 $0 \leq t' < \bar{t}_1$ 中, 包括在时点 0 上, 国家 2 都选择立刻放弃。所以, 在完全信息中, 均衡策略为国家 2 立刻选择放弃, 国家 1 得到全部利益。

这一完全信息均衡的特点是, 很高的解决方案和很高的观众成本, 都会导致更好的危机谈判结果。但是, 在这一均衡中, 危机永远不会发生, 因为弱小一方始终选择立刻投降。因此, Fearon 还构造了不完全信息模型, 其中, 每一方都不了解对方的解决方案。国家 i 的解决方案在区间 $[w_i, 0]$ 上的累积分布函数为 F_i 。

和完全信息博弈中的情况一样, 均衡表现为一个时点, 在这个时点之后, 没有任何一个国家选择放弃。Fearon 称这一时点为这一危机博弈的时域。就是说, 这一时域点被定义为 t_h , 对 $i = 1, 2$, 是 $Q_i(t)$ 在 $t > t_h$ 上不递增的最早时点。

Fearon 指出, 在任何的 PBE 中以下特征肯定成立:^②

(1) 两个国家都以 0 的概率立刻放弃。假设国家 1 在时点 t' 以正的概率密度放弃。国家 2 在放弃前, 显然有动机等待至少至时点 $t' + \varepsilon$ (ε 为某个很小

① 在这里的策略集中, 没有“永不放弃”这一策略。这就排除了如下的无意义的均衡: 双方都选择这一策略, 于是危机永远持续下去。

② 对这些结论的正式陈述和它们的证明, 见 Fearon(1994)。

的正数);这样做能使它以其观众成本上的微小增加,赢得 v 的概率显著上升。同样地,不存在在某个国家选择放弃的同时,对方发起进攻的均衡。如果国家 2 认为国家 1 会在时点 t' 以正的概率密度放弃,它就会推迟进攻直至 $t' + \varepsilon$ 。这就是说,两个国家都不可能计划在时点 t_h 选择放弃。

(2) 如果 $Q_j(t)$ 在 t' 上递增,国家 i 就不会在时点 t' 发起进攻。因为进攻决策产生的期望效用为负,所以,等待更长的时间对国家 i 有利,因为在这段时间中,对手可能选择放弃。

(3) 在任意地接近 t_h 的时间区间中,两个国家都以正的概率选择放弃。根据 t_h 的定义,至少一个国家必须在 t_h 之前的任意小的区间中,以正的概率选择放弃。假设这一点对国家 i 成立。再假设国家 j 在时点 $t' < t_h$ 之后,以 0 的概率选择放弃。根据结论(2),在 t' 和 t_h 之间国家 2 不会发起进攻。因此,国家 1 如果在时点 t' 之后以正的概率选择放弃,只会增加自己的观众成本——它知道自己将先于国家 2 选择放弃,但却依然在等待。

(4) 在时点 $t < t_h$, 两个国家都以 0 的概率选择进攻。这一结论直接得自结论(2)和结论(3)。因此,策略 $\{t > t_h, \text{进攻}\}$ 带给国家 i 的效用是:

$$U_i^q(t, w_i) = Q_j(t_h)v + (1 - Q_j(t_h))w_i \quad (10.8)$$

策略 $\{t < t_h, \text{放弃}\}$ 带给它的效用是:

$$U_i^q(t) = Q_j(t)v - (1 - Q_j(t_h))a_i(t) \quad (10.9)$$

由于 $U_i^q(t)$ 并不决定于 w_i , 所以,如果对所有的 t , $U_i^q(t)$ 为某个常数 k_i , $\{t < t_h, \text{放弃}\}$ 就是唯一的最优反应。假设并非如此。假设对某个 $t' < t_h$ 和对所有的 $t < t_h$, $U_i^q(t') > U_i^q(t)$ 。则使 $U_i^q(t') > U_i^q(t_h, w_i)$ 的所有类型,在时点 t' 选择放弃。国家 j 的最优反应于是为 $\{t > t', \text{放弃}\}$, 从而使 $U_i^q(t') = -a_i(t')$, 这显然与对所有的 $t < t_h$, $U_i^q(t') > U_i^q(t)$ 的前提相矛盾。

在这些结论的基础上,下面的引理能帮我们找到这一博弈的精炼贝叶斯均衡。

引理 10.4 在任何一个均衡中,如果两个国家都以正的概率选择持续冲突,则肯定存在有限时域 t_h 。

这一引理的逻辑相当直接。假设这一引理不成立,则存在某个 PBE, $Q_i(t)$ 在所有的 t 上递增。根据结论(2),国家 j 永远不会进攻。于是,根据结论(4), $U_i^q(t)$ 在所有的 t 上是常数。这就意味着:

$$Q_j(t) = \frac{k_i + a_i(t)}{v + a_i(t)}$$

但是,由于 j 永远不进攻,所以, $\lim_{t \rightarrow \infty} Q_j(t) = 1$ 。只有在 $k_i = v$ 或者 $\lim_{t \rightarrow \infty} a_i(t) = \infty$ 时,这一点才成立。如果 $k_i = v$, 则 $Q_j(0) = 1$ 意味着 j 肯定不会持续冲突。如果 $\lim_{t \rightarrow \infty} a_i(t) = \infty$, 则对任意大的 t , i 不会选择 $\{t, \text{放弃}\}$ 。

引理 10.5 在任何均衡中,如果时域为 t_h 并且冲突可能持续,(1)如果国家 i 选择 $\{t, \text{进攻}\}$, 则肯定有 $t \geq t_h$; (2)在 $t \geq t_h$ 时,当 $w_i > -a_i(t_h)$ 且仅当 $w_i \geq -a_i(t_h)$ 时,国家 i 选择 $\{t, \text{进攻}\}$ 。

这一引理的第一部分直接得自结论(2)和结论(3)。第二部分则得自这样一个事实:如果不考虑一些技术性因素,则当且仅当 $w_i \geq -a_i(t_h)$, $U_i^a(t, w_i) \geq U_i^q(t)$ 。①

根据引理 10.5,国家 j 在时点 t_h 发动进攻的先验概率是 $1 - F_j(-a_j(t_h))$ 。因此,国家 i 如果使危机持续至 t_h 后才选择放弃,实现的期望效用为:

$$u_i(t_h) = F_j(-a_j(t_h))v - (1 - F_j(-a_j(t_h)))a_i(t_h)$$

设 t_i^* 使得 $u_i(t_i^*) = 0$ 。就是说, t_i^* 的特征是,国家 i 在使危机持续至时点 t_i^* 和立刻放弃两种选择之间无差异。②

命题 10.1 设 t_i^* 是 $u_i(t_i^*) = 0$ 的唯一解,定义 $t^* = \min\{t_1^*, t_2^*\}$ 。在危机以正的概率持续的任何均衡中,时域肯定为 t^* 。

如果 t_h 大于 t^* , t_i^* 的值较小的国家有动机在时点 t_d 之前以 1 的概率放弃。这就与结论(3)矛盾。如果 t_d 大于 t^* , 则两个国家都有动机在放弃之前,在时点 t_d 上虚张声势一番。这与 t_d 的定义相矛盾。

这些结论直接产生了以下命题。

命题 10.2 重新命名两个国家从而使 $t^* = t_2^* < t_1^*$ 。设 $k_1 = u_1(t^*) > 0$ 。我们把国家 $i = 1, 2$ 的均衡策略表示为类型 w_i 的函数:

在 $w_i \geq -a_i(t^*)$ 时,在任何时点 $t > t^*$ 上,国家 i 使用策略 $\{t, \text{进攻}\}$ 。

在 $w_i < -a_i(t^*)$ 时,国家 i 使用策略 $\{t, \text{放弃}\}$ 。其中使用的任意纯策略产生如下累积分布:

$$\Psi_1(t) = \frac{1}{F_1(-a_1(t^*))} \frac{a_2(t)}{v + a_2(t)}$$

$$\Psi_2(t) = \frac{1}{F_2(-a_2(t^*))} \frac{k_1 + a_1(t)}{v + a_1(t)}$$

在时点 $t \leq t^*$, 国家 i 认为,国家 j 不会放弃的概率为:

$$\Pr(w_j \geq -a_j(t^*) | t) = \frac{v + a_i(t)}{v + a_i(t^*)}$$

在时点 $t > t^*$, 国家 i 的推断遵循贝叶斯法则,与对手选择进攻的策略相一致。对不在均衡路径上的任何的 $t > t^*$, i 认为 $w_j > -a_j(t^*)$, 即 w_j 依 F_j 分布且在 $-a_j(t^*)$ 截止。

证明: 设 $Q_i(t)$ 是国家 i 在时点 t 之前选择放弃的无条件概率。根据引理 10.5 和命题 10.1, $Q_i(t^*) = F_i(-a_i(t^*))$ 。国家 i 从 $\{t, \text{放弃}\}$ 中得到的效用因此为:

$$Q_j(t)v - (1 - Q_j(t))a_i(t)$$

① 在这些技术性细节中,包括排除掉 $Q_i(t)$ 有密度点的情况。

② 我们实际上是在假定, $a_i(t)$ 的值域足够大,从而使 $u_i(t_i^*) = 0$ 有唯一解。

为保证在任何时点 $t < t^*$ 上, 国家 i 在放弃与使危机持续之间无差异, 要求:

$$Q_j(t)v - (1 - Q_j(t))a_i(t) = u_i(t^*)$$

即

$$Q_j(t) = \frac{u_i(t^*) + a_i(t)}{v + a_i(t)}$$

由于只有类型 $w_j < -a_j(t^*)$ 会放弃, 所以这些类型必须以:

$$\frac{Q_j(t)}{F_j(-a_j(t^*))}$$

的速度放弃, 因此:

$$\Psi_j(t) = \frac{1}{F_j(-a_j(t^*))} \frac{u_i(t^*) + a_i(t)}{v + a_i(t)}$$

我们知道, 到时点 t , 在类型落在区间 $[w_j, a_j(t^*)]$ 的国家中, 有 $\Psi_j(t)$ 部分已经放弃, 所以:

$$\begin{aligned} \Pr(w_j \geq -a_j(t^*) | t) &= \frac{1 - F_j(-a_j(t))}{1 - F_j(-a_j(t)) + F_j(-a_j(t))(1 - \Psi_j(t))} \\ &= \frac{1 - F_j(-a_j(t))}{1 - Q_j(t)} \\ &= \frac{(1 - F_j(-a_j(t)))(v + a_i(t))}{v - u_i(t^*)} \\ &= \frac{v + a_i(t)}{v + a_i(t^*)} \end{aligned} \quad \square$$

在这个 PBE 中, 有低的解决值的国家 i 所选择的放弃速度, 使得有低的解决值的国家 j 在放弃与持续至时点 t^* 之间无差异。在时点 t^* , 两个国家都发动进攻, 因为它们知道, 所有有低的解决值的对方国家都已经选择了放弃, 任由危机持续下去只会导致更大的观众成本。

有低的解决值的国家 j 所选择的放弃速度, 为什么使得同样有低的解决值的国家 i , 在放弃和持续冲突之间无差异呢? 如果国家 j 以更快的速度放弃, 则所有有低的解决值的国家 i 就会断定在每个时点 t , 国家 j 更可能为强壮型。这就导致有低的解决值的国家 i 亦更快地放弃。于是, 国家 j 就会断定, 国家 i 的类型更强壮, 于是更快地放弃。在极限点, 所有有着低的解决值的国家都会在时点 $t = 0$ 放弃[见结论(1)]。

再来探讨有着低的解决值的国家 i , 以比均衡更慢的速度放弃可能会导致的后果。在每个时点 t , 国家 j 会推定剩下的类型要弱于相应的均衡类型。这就导致国家 j 也以更慢的速度放弃, 进而导致国家 i 持续更长时间, 如此等等。这种动态性导致所有类型的国家愿意维持至时点 t^* 。所以, 这不可能是均衡; 在时点 t^* 之前放弃和在时点 t^* 发动进攻之间, 有着低的解决值的国家显然偏好前者。

这个模型的一个重要结论是,能够发生很大的观众成本是件有好处的事。高成本国家能够更好地使对手相信它被“套牢”在冲突中(就是说,有更多的类型愿意持续至时域点,然后开始进攻)。而这与一个简单的直观认识——那些从更早的放弃中失去的利益最多的国家会更早地投降——相矛盾。一旦危机发生,信号价值就压倒了这一效应。通过在模型中引入危机发生阶段, Fearon 认为,这一分析框架就能解释为什么民主国家(即对观众成本高度敏感的国家),不大可能挑起冲突,但是却更可能在冲突中取胜。

10.8 习题

习题 10.1 对 $i \in \{A, B\}$, 设 $u_i(x_i) = \ln x_i$ 。求解 Nash 讨价还价解, 将其表示为分歧值的函数。

习题 10.2 证明: 满足四个 Nash 公理的唯一的一个讨价还价解是 Nash 解。

习题 10.3 在 Rubinstein 讨价还价模型中, 设 $\delta_1 = \delta_2$ 且 $d_1 = d_2 = 0$, 并假定 $u_i(x_i) = x_i^\alpha$, 其中 $0 < \alpha < 1$ 。计算 SPNE 份额。风险厌恶(就是说, 更低的 α 值)有怎样的影响?

习题 10.4 设在封闭规则下的 Baron-Ferejohn 模型中, $u_i(x_i) = x_i^\alpha$ 。证明初始提案者的份额在 α 上递减。

习题 10.5 在第 10.3.2 节的模型中, 设 $N = 3$, $m = 1$, $\frac{1-\delta}{3-2\delta} < q \leq \frac{1}{3}$ 。计算在其中 $v_A = v_B$ 的混合策略均衡。

习题 10.6 在第 10.3.2 节的模型中, 对一般的 N 和 m 值, 计算 v_A 和 v_B 。

习题 10.7 在第 10.3.3 节的模型中, 假设有两个团体在讨价还价, 称这两个团体分别为 1 和 2。设 m_i 是第 i 个团体的成员数量, $m_1 + m_2 = N$ 。假设团体 i 的所有成员都有否决权, 就是说, 如果他们反对某项提案, 则需要有 k_i 的票数才能够推翻他们的否决, 其中, $0 < k_i < N$ 。假设 $k_1 > k_2$ 。对每个团体的成员, 计算持续值。

习题 10.8 在第 10.5.1 节的模型中, 假设有两轮讨价还价: 如果 A 拒绝了 B 的初始方案, A 提出自己的方案。如果双方在第二轮中达成一致, 则收益以因子 δ 贴现。这一博弈的 PBE 是什么? 假设在 B 的分歧值上存在着不完全信息, $u_B = 0$ 的概率为 π , $u_B = d$ 的概率为 $1 - \pi$ 。构造这一博弈的一个 PBE。它是这一模型唯一的 PBE 吗?

习题 10.9 证明: 在图 10.10.3 中描述的 Matthews 模型中, 不存在所有的总统类型发表相同讲话的全域至圣均衡。

习题 10.10 在第 10.7.1 节的模型中, 假设 c_1 在区间 $[\underline{c}, \bar{c}]$ 上均匀分布。求精炼贝叶斯均衡。假设在讨价还价前, 存在“空谈”。证明: 如果 $\bar{c}$ 足够大, 则存在 PBE, 高成本类型的国家 1, 向国家 2 揭示关于其成本的信息。

机制设计和代理理论

此前讨论分析技术，专用于分析策略性行为主体在具体博弈中是如何行动的。在某些社会环境中，当博弈规则十分明确时，这一方法能为我们提供直觉认识和经验预测。还有一种理论探讨与此略有不同的问题：给定想要的结果，策略性行为主体之间的何种博弈能产生这种结果？这样的博弈存在吗？

博弈论中探讨此类问题的领域，称为机制设计。在这种分析框架中，机制设计者或者委托人，选择一个贝叶斯博弈或者一个机制，由代理人进行博弈。机制设计的例子包括，设计税制——借此诱使代理人揭示自己为公共项目提供资金的意愿；设计拍卖——借此最大化销售收入；选民设计选举函数——借此激励政府官员恪尽职守。

一般来说，对机制的选择是个最大化问题：机制设计者选择一个博弈，借此在代理人理性地进行这一博弈的条件下，最大化设计者的效用。如果设计者的偏好是个社会偏好，机制设计就是个规范问题。这种规范问题的典型例子，是通过决定公共品生产的规则选择，最大化个人效用之和。从这种规范诠释视角看，机制设计与社会选择理论有密切联系。有一种机制设计理论被称为执行（implementation）理论，它试图找到选择函数（从代理人类型到集体决策的映射）和实现这一选择函数的机制。这类选择函数被称为是可执行的（implementable）选择函数。Gibbard-Satterthwaite 定理便属于这一理论范畴。

机制设计理论的各种应用虽然常常是规范性的，但是，我们可以用它产生实证结论。例如，机制设计常用于研究委托代理关系，目的是探讨缺乏信息的委托人如立法机构或政府首脑，是否能够诱使拥有信息的代理人如委员会或者政府官员，为委托人的利益行动。

在大部分经济学应用中，有足够丰富的博弈类型可供设计者从中做出选择。他通常能够设计出有着精巧的奖惩机制的合同。这种假定在经济背景中有其合理性，因为在经济问题中，第三方如法庭等，能强制执行各种复杂合同，并且巨额货币奖惩是合法手段。但是，在政治背景中，我们不能假定委托人能够推行这种奖励机制，因为在政治背景中，常常不存在来自第三方的强制力；并且，货币激励常常被法律禁止或被社会禁止。所以，在介绍机制设计的基本概念和结论后，我们探讨委托人受到更多约束的激励问题。

11.1 例子

我们探讨一个典型例子。设在政治科学系中,有 n 个教师和一个系主任。系主任在考虑是否购买一台咖啡机。咖啡机的价格为 c ,系主任想知道系里教师对这台咖啡机的价值认定是否足够高,从而值得发生这笔开支。每个教师都将这台机器的价值认定为 $\theta_i \in \mathbb{R}_+$ 。系主任是个边沁主义者但不喝咖啡^①,当且仅当 $\sum_{i=1}^n \theta_i \geq c$ 时,他才会购买这台机器。系主任不知道各个教师对咖啡机的价值认定,但知道每个教师的类型服从概率分布 $F(\cdot)$ 。

系主任应该怎么做呢?一种办法是私下向每个成员询问对咖啡机的价值认定,如果所揭示的总价值认定超过价格 c ,他就用系里经费购买咖啡机。问题是,一些教师知道,夸大自己的价值认定,借此提高购买概率,对自己有利。^②就是说,这种办法诱生了一个博弈,其中每个教师的最优反应是报告高于真实值的价值认定。所以,对于确定这台机器是否值得购买而言,这种办法没有起作用。更好的解决办法是要求每个教师按照自己的价值认定出资,如果筹措的资金总额超过 c ,就购买这台机器,剩下的钱购买咖啡豆。如果资金总额没有达到 c ,系主任把钱退回给每个人。但是,这种机制也有缺陷。善于计谋的学者会低报自己的价值认定,希望借此搭便车,占同事便宜。这一办法导致了集体行动问题。

这两种办法都不能很好地起作用,但是,我们可以用机制设计理论找到一些特别简单的机制,来了解教师的偏好。Groves(1973)和 Clarke(1971)指出,下面的机制有着我们想要的特征:

- 每位教师用电子邮件把自己的价值认定 m_i 告知系主任。
- 如果 $\sum_{i=1}^n m_i \geq c$, 就购买这台咖啡机;否则,就不购买。
- 如果购买了咖啡机,就向教师 i 收取数量为 $t_i(m_i, m_{-i}) = c - \sum_{j \neq i} m_j$ 的资金。
- 如果没有购买咖啡机,就不收钱。

在这一机制下,每个教师都有动机揭示自己真实的价值认定,而不会顾及其他人的价值认定。这里的一个关键特征是,教师 i 发出的信息只是通过它对是否购买咖啡机的最终决策的影响,间接地影响自己的出资额。每个教师的出资额决定于其他所有人发出的信息。

所有教师都有动机提供真实信息。为说明这一点,我们分析类型为 θ_i 的第 i 个教师的决策。假设他没有说实话,即他宣布 m_i , $m_i < \theta_i$ 。他低报自己的类型的行为,只在改变购买机器的概率时,或者说在:

① Jeremy Bentham(1748~1832)是英国著名哲学家,他所倡导的道德体系以“为最多的人谋取最大的福利”为基础。历史记载表明,他喝咖啡。

② 这里假定教师并不关心系里的预算经费是否有其他用途。

$$\sum_{j \neq i} m_j + m_i < c < \sum_{j \neq i} m_j + \theta_i$$

时才影响最终结果。根据上面的条件, $c - \sum_{j \neq i} m_j < \theta_i$, 就是说, 即在报告了自己真实的价值认定时所要缴纳的出资额 $c - \sum_{j \neq i} m_j$, 小于他对咖啡机的价值认定。这个教师没有报告自己的真实类型的行为, 只会使他的效用下降。

再看教师是否有动机夸大自己对咖啡机的需求。假设他发出信息 m_i , $m_i > \theta_i$ 。只有当这一夸大的信息导致系里购买这台机器而真实信息导致不购买这台机器时, 这种歪曲真相的行为才影响 i 的效用。这实际上要求:

$$\sum_{j \neq i} m_j + m_i > c > \sum_{j \neq i} m_j + \theta_i$$

即 $c - \sum_{j \neq i} m_j > \theta_i$ ——教师 i 的出资额高于他对咖啡机的价值认定。所以, 撒谎没有给他带来好处。

Groves-Clarke 机制除了能促进系里事务上的诚实性外, 还有另一个相当好的特征: 当且仅当全体教师对咖啡机总的价值认定超过它的价格时, 才会购买这台咖啡机。但是, 它有一个缺陷: 它没有“平衡预算”。如果这台机器最终被购买, 系主任收上来的教师们的出资总额为:

$$\sum_{i=1}^n t_i(m_i, m_{-i}) = nc - \sum_{i=1}^n \sum_{j \neq i} m_j$$

它大于或等于 c 。当然了, 在系主任看来, 这并不是个问题, 而只是对他收发和阅读这些邮件的一点补偿。

11.2 机制设计问题

我们现在更抽象地探讨机制设计问题。 n 个代理人和 1 个机制设计者(称为第 0 个人)构成集合 N 。设计者最终选择政策 $x \in X$ 。每个代理人有类型 $\theta_i \in \Theta$, 这是他的私人信息。类型向量 θ 服从联合分布函数 $F(\theta)$ 。代理人有贝努里效用函数 $u_i: X \times \Theta^n \rightarrow \mathbb{R}$, 就是说, 效用决定于设计者所选择的政策和代理人的类型。机制设计者有贝努里效用函数 $u_0(x, \theta)$ 。在许多情况中, 代理人只关心自己的类型, 但是在这里, 代理人的收益是类型组合的函数。机制设计问题因此被定义为 $\langle \Theta, F(\cdot), X, u \rangle$ 。

在一般应用中, 机制设计者诱使代理人发出信号向量。设计者为每个代理人选择一个信号空间 M_i ; 并选择政策函数 $p: \prod_{i \in N} M_i \rightarrow X$, 它为每个可能的信号组合 m 规定一个政策, 其中, $m = (m_1, m_2, \dots, m_n) \in M = \prod_{i \in N} M_i$ 。所以, 机制可被定义为 $\langle M, p(\cdot) \rangle$ 。

给定信号空间和政策函数, n 个代理人进行贝叶斯标准式博弈, 其中策略集

$S_i = M_i$, 收益函数为 $u_i(p(m), \theta)$, 即效用函数 u 和政策函数 p 的复合。

前述咖啡机购买机制适用于这一框架。系主任是设计者, 他选择信号空间 $M = \Theta$, 并执行政策“如果 $\sum_{i=1}^n m_i \geq c$ 就购买并向代理人 i 收取 $c - \sum_{j \neq i} m_j$ ”。全体教师然后进行贝叶斯标准式博弈。

在给定的机制下, 要确定代理人的行为, 还需要明确理性所具有的形式。一种办法是, 只当代理人在由这一机制所诱生的博弈中有占优策略时, 才做出预测。这种机制被称为可用占优求解的机制 (dominance solvable mechanism)。代理人也可能使用贝叶斯 Nash 均衡。所设计的机制多好地起作用的问题, 显然决定于代理人如何进行所诱生的博弈。

在本章中, 我们探讨贝叶斯 Nash 均衡。在经济学中, 有很大一块文献使用其他的解概念。例如, Groves-Clark 机制源自关于占优策略上的执行的文献。由于我们只探讨贝叶斯 Nash 均衡, 我们要求选择函数 $g: \Theta \rightarrow X$ 满足以下条件: 存在机制 $(M, p(\cdot))$, 如果代理人使用贝叶斯 Nash 均衡, 则最终结果为选择函数所选择的政策, 就是说, 对每个 $\theta \in \Theta$, $p(m(\theta)) = g(\theta)$ 。

从机制设计者角度看, 所设计的机制必须能够执行某个选择函数。设计者如果想在贝叶斯 Nash 策略上执行函数 $g(\cdot)$, 他对机制 $(M, p(\cdot))$ 的设计, 必须使得所生成的博弈有贝叶斯 Nash 均衡, 并且在这一均衡, 代理人使用策略 $m_i^*(\cdot)$, 从而使得 $p(m_1^*(\theta_1), m_2^*(\theta_2), \dots, m_n^*(\theta_n)) = g(\theta)$ 。这就是说, 在设计机制时, 必须了解对所设计的机制具有的激励。设计者知道在所设计的机制中代理人使用均衡策略, 从而知道这些激励如何影响代理人的行为。

定义 11.1 给定机制设计问题 $(\Theta, F(\cdot), X, u)$, 选择函数 $g(\cdot)$ 在贝叶斯 Nash 策略上是可执行的, 如果存在机制 $(M, p(\cdot))$, 它产生贝叶斯 Nash 均衡 $m_i^*(\cdot)$, 且对每个 $\theta \in \Theta$,

$$p(m_1^*(\theta_1), m_2^*(\theta_2), \dots, m_n^*(\theta_n)) = g(\theta)$$

一个博弈如果可用占优概念求解, 则存活下来的策略组合就是贝叶斯 Nash 均衡。这就意味着, 给定一个机制设计问题, 在占优策略上可执行的选择函数集, 是在贝叶斯 Nash 策略上可执行的选择函数集的子集。我们已经定义了可执行性, 随之而来的是个更难的问题。由于可供选择的机制数量很大, 如何知道哪个选择函数是可执行的呢?

我们的第一个结论是揭示原理 (revelation principle), 它使我们只需探讨直接机制集这个更小的集合, 从而简化找到可执行的选择函数的过程。在直接机制 (direct mechanism) 中, 代理人被要求直接报告自己的类型。就是说, 在直接机制中, $M_i = \Theta_i$ 。揭示原理指出, 如果存在执行选择函数 $g(\cdot)$ 的机制, 就肯定存在执行 $g(\cdot)$ 的直接机制。这一有力的结论使我们不必考虑所有可能的机制, 而只需探讨直接机制。

揭示原理虽是个一般性原理, 但是它的证明却十分简单。

定理 11.1 (贝叶斯 Nash 策略上的揭示原理) 给定机制设计问题 $\langle \Theta, F(\cdot), X, u \rangle$, 如果选择函数 $g(\cdot)$ 在贝叶斯 Nash 策略上是可执行的, 则存在一个直接机制 $\langle \Theta, p(\cdot) \rangle$, 它在贝叶斯 Nash 策略上执行 $g(\cdot)$ 。

证明: 假设存在非直接机制 $\langle M, p'(\cdot) \rangle$, 它在贝叶斯 Nash 策略上执行 $g(\cdot)$ 。我们用这一机制构造一个执行选择函数 $g(\cdot)$ 的直接机制。设在机制 $\langle M, p'(\cdot) \rangle$ 所诱生的博弈中, 在其中的一个贝叶斯 Nash 均衡上, 参与者 i 使用策略 $s_i(\cdot)$ 。我们构造一个直接机制: 代理人被要求发出信息 $m_i \in \Theta_i$, 然后对每个 $\theta \in \Theta$, 通过函数 $p(\theta) = p'(s_1(\theta_1), \dots, s_i(\theta_i), \dots, s_n(\theta_n))$ 选择政策。我们只需验证在这一直接机制中, 发出真实信息即 $m_i(\theta_i) = \theta_i$, 构成贝叶斯 Nash 均衡。首先, 假设所有代理人 $N \setminus \{i\}$ 都使用这一说真话策略。代理人 i 如果也使用说真话策略, 则对每个 θ , 最终结果为 $g(\theta)$ 。现在假设在直接机制中, 对类型为 $\theta'_i \in \Theta_i$ 的代理人 i , 存在着更有利的选择 m'_i , $m'_i \neq \theta'_i$ 。在直接机制中, 如果代理人 i 能选择 m'_i , 则 $m'_i \in \Theta_i$ 。但是, 由于在 $\langle M, p'(\cdot) \rangle$ 所诱生的博弈中, $s_i(\theta_i): \Theta_i \rightarrow M_i$, 所以肯定存在 $s_i(m') \in M_i$ 。这意味着

$$\int_{\theta_{-i} \in \Theta_{-i}} u_i(p'(s_1(\theta_1), \dots, s_i(m'), \dots, s_n(\theta_n))) dF(\theta_{-i} | \theta_i) > \int_{\theta_{-i} \in \Theta_{-i}} u_i(p'(s_1(\theta_1), \dots, s_i(\theta'_i), \dots, s_n(\theta_n))) dF(\theta_{-i} | \theta_i)$$

这一表达式与在由 $\langle M, p'(\cdot) \rangle$ 所诱生的贝叶斯博弈的均衡中 $s_i(\cdot)$ 是最优反应的事实相矛盾。于是, 定理得证。 $\square$

有一点需要指出。这一定理只是说, 存在一个机制, 在贝叶斯 Nash 均衡中, 它执行某个选择函数。还可能存在其他均衡, 导致不同的集体选择。建议读者参阅 Palfrey 和 Srivastava(1989), 它探讨了在要求机制有唯一均衡时的机制设计问题。

我们可以把揭示原理推广至很多均衡概念上。如前所述, 一个重要例子是占优策略上的执行。习题 3 给出了它的定义, 并请读者证明这种情况中的揭示原理。其中的道理与前述证明类似。假设在某个机制中, 参与者有占优策略。则在恰当定义的直接机制中, 如果它能为诚实信号提供原始机制中的收益, 说真话便是占优策略。

有了揭示原理, 我们只需探讨说真话的直接机制, 就能够回答选择函数是否可执行的问题。某个选择函数如果无法在这种机制中被执行, 则在任何机制中, 都无法被执行。

在继续理论探讨之前, 我们看几个例子。

11.3 应用: 民调

假设有 n 个选民 (n 为奇数), 即 $N = \{1, 2, \dots, n\}$, 每个选民在 $\mathbb{R}$ 上有对称单峰偏好。机制设计者不清楚选民的理想点, 但可以向每个选民询问, 然后选择政策

$x \in \mathbb{R}$ 。假设每个理想点 θ_i 服从 $\mathbb{R}$ 上的分布 $F(\cdot)$ 。一个很自然的问题是, 是否存在某些机制, 能诱使选民在占优策略中揭示自己的理想点。Herve Moulin (1980) 指出, 答案是肯定的。我们来分析这样一种机制。它要求每个选民报告自己的理想点 $m_i \in \mathbb{R}$, 然后选择政策, 使所选择的政策等于所有选民报告的理想点的中位值, 即 $x(m) = \text{median}(m)$ 。在这里, 说真话是最优反应。我们来验证这一结论。设 $x(m_i, m_{-i})$ 为信号组合 m 的中位数, 这一信号组合包含 m_i 和其他 $n-1$ 个选民所报告的类型; 设 $\underline{x}(m_{-i})$ 和 $\bar{x}(m_{-i})$ 分别为 $\text{median}(m_{-i})$ 的下界和上界。^①

假设选民 i 报告自己的理想点为 y_i 。如果 $y_i < \underline{x}(m_{-i})$, 则 $\underline{x}(m_{-i})$ 为所有选民所报告的理想点的中位值, 即 $x(y_i, m_{-i}) = \underline{x}(m_{-i})$ 。如果 $y_i > \bar{x}(m_{-i})$, 则 $x(y_i, m_{-i}) = \bar{x}(m_{-i})$ 。如果 $y_i \in \{\underline{x}(m_{-i}), \bar{x}(m_{-i})\}$, 则 y_i 就是中位值, 即 $x(y_i, m_{-i}) = y_i$ 。因此, 报告自己的理想点为 y_i , 只会导致三种结果中的一种, $\underline{x}(m_{-i})$ 、 $\bar{x}(m_{-i})$ 和 $y_i \in \{\underline{x}(m_{-i}), \bar{x}(m_{-i})\}$ 。

如果 $\theta_i \in \{\underline{x}(m_{-i}), \bar{x}(m_{-i})\}$, 则报告 y_i 不可能带来好处, 因为如果报告 $m_i = \theta_i$, 选民 i 能得到自己的理想点。假设 $\theta_i < \underline{x}(m_{-i})$ 。选民 i 现在的最好的可能结果是 $\underline{x}(m_{-i})$ 。只要所报告的理想点小于 $\underline{x}(m_{-i})$, 包括报告 θ_i , 都能实现这一结果。假设 $\theta_i > \bar{x}(m_{-i})$, 则选民 i 报告 $m_i = \theta_i$ 能够弱最大化自己的效用。所以说, 不管其他选民如何报告自己的理想点, 对选民 i 来说, 最优策略是报告 $m_i = \theta_i$ 。

在 Moulin 所探讨的机制中, 选民被要求报告一个数字 (其理想点)。他要求这一机制满足两个条件: 匿名性和效率。第一个条件即匿名性, 只是要求这一机制以完全相同的方式对待每个人。

定义 11.2 机制 $g: \mathbb{R}^n \rightarrow X$ 是匿名的, 如果对每个排列 $\pi: N \rightarrow N$,

$$g(m_1, \dots, m_i, \dots, m_n) = g(m_{\pi(1)}, \dots, m_{\pi(i)}, \dots, m_{\pi(n)})$$

Moulin 的效率条件基本上是 Arrow 的帕累托条件。

定义 11.3 机制 $g: \mathbb{R}^n \rightarrow X$ 是有效率的, 如果对每种类型组合 $y = (y_1, \dots, y_n)$, $g(y)$ 具有帕累托效率 (就是说, 不存在其他某个政策 $z \neq g(y)$ 使得在 z 和 $g(y)$ 之间, 每个人都弱偏好 z , 并且有某人严格偏好 z)。

Moulin 证明了以下结论:

定理 11.2 如果偏好是单峰的, 则每个有效率的和匿名的策略防范机制具有以下形式: 在所报告的数字 $m_1, m_2, \dots, m_n$ 上, 增加 k 个固定数字 $a_1, \dots, a_k$, 然后选择这一列数字的中位值。

我们已经讨论过这一机制为什么能防范策略性行为, 这里把 Moulin 定理的证明留做习题。Meirowitz (2004a) 分析了在候选人利用民调知会选民自己的政策选择时, 在民调中存在的信息揭示激励。他发现, 受访者通常有动机说假话。我们设计了一道习题, 请读者分析 Meirowitz 的结论。

① 某个实数集如果包含偶数个不同元素, 则有两个中位数。

11.4 拍卖理论

从出售无线波段到在 eBay 上销售商品,拍卖几乎被用于一切商品的销售。由于它的这种作用,经济学家们就如何最优地设计拍卖机制,从而产生最大销售收入和实现配置效率,发展形成了庞大的理论体系。政治科学家们一般不关心拍卖机制的应用,但是,他们应该重视拍卖理论的一些基本原理。原因有几点。首先,一些重要的政治现象可以看作是某些类型的拍卖。其次,拍卖理论向我们展示了如何用机制设计理论研究对制度的选择。

政治科学实际上没有忽视拍卖理论。它的一个经典应用便是政治科学中的“收买”(favor buying)模型:相互竞争的利益团体向政治家提供贿赂,以求获得政府合同或得到政策优惠。最近,本书作者之一就把竞选作为一种拍卖形式进行研究(Meirowitz, 2004b)。

在标准的拍卖问题中,有一个卖方(第 0 个人)和 n 个潜在的买方,即 $N = \{1, \dots, n\}$ 。每个买方赋予所拍卖商品的价值为 $\theta_i \in \mathbb{R}_+$ 。这种价值认定是每个买方的私人信息。先验推断是,买方的价值认定服从独立同分布,分布函数 $F(\cdot)$ 二阶可微。在拍卖模型中,支付价格 b_i 而赢得商品且对商品的价值认定为 θ_i 的买方 i ,得到效用 $\theta_i - b_i$ 。没有拍得商品的买家,如果没有发生任何支付,收益为 0。我们首先探讨几个得到普遍研究的拍卖机制。

11.4.1 次高价格拍卖和价格升序拍卖

我们探讨两种拍卖机制。在次高价格拍卖中,每个参与者递交封闭报价(sealed bid) b_i ,出价最高的人得到商品。但是,他所支付的金额为次高报价,其他人不发生任何支出。在价格升序拍卖中,所有买家首先举起自己的号码牌,拍卖师按照升序喊出价格,10 元,11 元,12 元……如果在某个价格上,除一人外,所有人手中的牌子都放下,这件商品就卖给这个在喊出的价格上依然举牌的人。

这两种拍卖方式都被广泛使用,在两种拍卖方式中的竞拍过程似乎差异甚大。但是,从博弈论角度看,这两种机制完全相同。首先要指出的是,它们生成不同的博弈(一个为参与者同时行动的贝叶斯博弈,另一个是参与者序贯行动的贝叶斯博弈)。但是,它们中的激励机制相同。我们用揭示原理来验证这一结论。设想每个买方把自己对商品的价值认定 θ_i 输入到电脑中,然后由电脑完成两件事:(1)它把这一价值认定提交给次高价格拍卖(就是说,它使用策略 $b_i = \theta_i$)。(2)它让一个机器人在升序拍卖中使用如下策略:举牌,直至所宣布的价格超过价值认定 θ_i 。在这两种情况中,对商品有最高价值认定的买家赢得商品,所支付的价格都等于对商品的次高价格认定。

在次高价格拍卖中,报价 $b_i = \theta_i$ 是占优策略。假设报价为 $b_i \neq \theta_i$ 。 $b_i^{\max} = \max_{j \neq i} b_j$

表示其他买家提交的最高报价。于是,有三种可能性: $b_i^{\max} = \theta_i$ 、 $b_i^{\max} < \theta_i$ 和 $b_i^{\max} > \theta_i$ 。在第一种情况中, $b_i = \theta_i$ 导致平局,期望效用为 0。使用策略 $b_i > \theta_i$, 则能在价格 $b_i^{\max} = \theta_i$ 上赢得商品,得到的效用为 0。使用策略 $b_i < \theta_i$, 得不到商品,效用为 0。在第二种情况下,使用策略 $b_i = \theta_i$, 赢得商品,支付 $b_i^{\max}$, 他的收益为 $\theta_i - b_i^{\max}$ 。使用策略 $b_i > \theta_i$, 赢得商品,支付 $b_i^{\max}$, 收益为 $\theta_i - b_i^{\max}$ 。使用策略 $b_i < \theta_i$, 得不到商品,收益为 0。在第三种情况下,使用策略 $b_i = \theta_i$, 得不到商品,收益为 0。使用策略 $b_i > \theta_i$, 或者能得到商品,支付 $b_i^{\max}$, 但他对商品的价值认定为 $\theta_i < b_i^{\max}$; 或者,得不到商品,收益为 0。因此说,使用策略 $b_i = \theta_i$, 至少和使用其他策略一样得好,而且能够防止事后后悔:在他愿意支付的价格上没有赢得商品,和在他不想支付的价格上赢得了商品。

类似地,在价格升序拍卖中,一旦价格超过 θ_i 就放下号码牌的策略是占优策略。在价格达到 θ_i 前,应该放下牌子吗? 这样做只会导致自己得不到商品,产生的收益为 0。所以,过早放弃不能带来更大好处。而且,如果在价格 $p' < \theta_i$ 上,其他人都放下了号码牌,则继续举牌的策略能使自己以价格 p' 得到商品,得到效用 $\theta_i - p'$ 。假设在价格超过 θ_i 以后,自己继续举牌。则或者得不到商品,得到效用 0, 或者得到了商品,但是所支付的价格超过他对商品的价格认定。

11.4.2 最高价格拍卖和价格降序拍卖*

在最高价格拍卖中,所有竞拍者同时提交自己的报价。出价最高的人赢得商品,并按自己的报价支付给卖方。在价格降序拍卖中,每个竞拍者有一个号码牌。拍卖师从某个高价位上开始喊价,并逐步降低价格,如 100 元,99 元,98 元……最早举起号码牌的竞拍者得到商品,支付当时喊出的价格。正如次高价格拍卖类似价格升序拍卖,这里的两种拍卖方式中所包含的同时行动博弈和扩展式博弈也十分类似。我们只探讨最高价格拍卖,而把构造价格降序拍卖中的精炼贝叶斯均衡留做习题。我们根据 Krishna(2002)不十分严格地推导最高价格拍卖的均衡策略。我们首先猜测均衡策略,然后证明这一猜测与贝叶斯 Nash 均衡相一致。

假设竞拍人 $j \neq i$ 的策略为可微的和严格递增的函数 $b(\cdot): \mathbb{R}_+ \rightarrow \mathbb{R}_+$ 。一位竞拍人如果为商品支付的价格等于其对商品的价值认定,或者他没有竞得商品,则期望效用为 0。我们用随机变量 $b_i^{\max}$ 表示其他竞拍人 $N \setminus \{i\}$ 在使用所猜想的这一策略时,给出的最高报价。由于得不到商品时的收益为 0,所以,类型为 θ_i 的竞拍人,从使用策略 $b_i = b(\theta_i)$ 中得到的期望效用为:

$$\Pr(b_i^{\max} < b_i)[\theta_i - b_i]$$

如果其他拍卖人使用 $b(\cdot)$, 则 $\Pr(b_i^{\max} < b_i)$ 就是从 $F(\cdot)$ 中 $n-1$ 次地抽取 θ_j , 每次得到的 θ_j 的值都小于 $b^{-1}(b_i)$ 的概率。由于类型为独立随机变量,所以,这一概率为 $F(b^{-1}(b_i))^{n-1}$ 。期望效用于是为:

$$F(b^{-1}(b_i))^{n-1}[\theta_i - b_i]$$

最优的 b_i 是一阶必要条件的解, 即:

$$(n-1)F(b^{-1}(b_i))^{n-2}f(b^{-1}(b_i))\frac{db^{-1}(b_i)}{db_i}[\theta_i - b_i] - F(b^{-1}(b_i))^{n-1} = 0$$

在所猜想的均衡中, $b(\theta_i) = b_i$, 于是:

$$(n-1)F(\theta)^{n-2}f(\theta)\theta = \frac{db(\theta)}{d\theta}F(\theta)^{n-1} + (n-1)F(\theta)^{n-2}f(\theta)b(\theta)$$

等号右边的项为 $\frac{d}{d\theta}(F(\theta)^{n-1}b(\theta))$ 。于是, 我们可以把一阶条件表述为:

$$\frac{d}{d\theta}(F(\theta)^{n-1}b(\theta)) = (n-1)F(\theta)^{n-2}f(\theta)\theta$$

利用微分第一定理 $\left[\int \frac{df}{dx} dx = f(x) \right]$, 我们把这个条件表述为:

$$b(\theta) = \frac{\int_0^\theta \theta' (n-1)F(\theta')^{n-2}f(\theta')d\theta'}{F(\theta)^{n-1}} \quad (11.1)$$

(11.1)式右边的项是在 $b^{\max} < \theta$ 时 $b^{\max}$ 的条件期望。要证明 $b(\cdot)$ 确实是均衡报价函数, 等同于确定(11.1)式所定义的局部极值是全局最大值。我们将这一证明留做习题。

11.4.3 收入等价原理*

最高价格拍卖和次高价格拍卖有不同的出价策略。于是, 我们自然要问, 卖方偏好哪种拍卖方式? 拍卖理论的一个基本结论被称为收入等价原理, 它回答了这个问题(Riley and Samuelson, 1981; Myerson, 1981)。这一原理指出, 当类型为独立变量时, 拍卖的期望收入只决定于卖方胜出的概率——这个概率是其类型 θ_i 的函数——和有最低的可能类型的卖方所实现的效用。这一原理的一个直接结论是, 如果价格认定相互独立, 上述所有拍卖方式给卖方产生的期望收入相同。本节证明这一结论。

首先回到本章开始时所定义的符号上。设 $M = \Theta$ 表示信号空间。 n 个参与者中的每一个人, 对所拍卖商品有自己的价值认定, 它们独立地抽取自 Θ 上的可微分布 $F(\cdot)$ 。这一分布的密度为 $f(\cdot)$ 。为方便分析, 设 Θ 为区间 $[0, k]$ 。集合 $\Delta(N)$ 包含竞拍人集合 N 上的所有概率分布。设机制 $w(m_1, \dots, m_n): M^n \rightarrow \Delta(N)$ 和 $t(m_1, \dots, m_n): M^n \rightarrow \mathbb{R}^n$ 对每个信号组合 m , 规定了赢家的身份上的概率分布和支付组合。于是, 在最高价格拍卖中, 赢家为:

$$w(m) = \begin{cases} \frac{1}{\#\{\text{argmax}\{m_j\}\}}, & \text{如果 } i \in \text{argmax}\{m_j\} \\ 0, & \text{在其他情况中} \end{cases}$$

支付为:

$$t_i(m) = \begin{cases} \frac{m_i}{\#\{\text{argmax}\{m_j\}\}}, & \text{如果 } i \in \text{argmax}\{m_j\} \\ 0, & \text{在其他情况中} \end{cases}$$

如果不考虑平局的可能性,则在 $i \in \text{argmax}\{m_j\}$ 时,这两项可分别简化为 1 和 m_i ; 否则,分别为 0。在后面将分析的全支付的最高价格拍卖中, $t_i(m) = m_i$ 。

收入等价原理所探讨的是,在拍卖中贝叶斯 Nash 均衡上的期望收入。标准的拍卖有前面所定义的“赢家通吃”映射 $g(m)$ 。给定递增报价函数 $b(\theta_i)$,期望收入为 $ER = \sum_{i=1}^n \int_0^k t_i(b(\theta_i)) dF(\theta_i)$ 。现在给出这一定理:

定理 11.3 设价值认定为独立同分布。在每个标准拍卖中,在每个对称的和递增的均衡中,如果价值认定为 0 的竞拍人的期望支付是 0,则它们有相同的期望收入值。

证明: 给定一个标准拍卖,其对称的和递增的均衡用函数 $b(\cdot)$ 定义。设 $t^+(\theta_i)$ 表示价值认定为 θ_i 的竞拍人的期望支付。假设竞拍人 i 的价值认定为 θ'_i ,我们分析其报价 $z = b(\theta'_i)$ 。在其他所有人使用策略 $b(\cdot)$ 时,报价 z 带给竞拍人 i 的期望效用为:

$$\theta'_i F(\theta'_i)^{n-1} - t^+(\theta'_i)$$

把这一期望效用对类型 θ'_i 求导,得到一阶条件:

$$(n-1)\theta'_i f(\theta'_i)^{n-2} = \frac{\partial t^+(z)}{\partial z} \Big|_{z=\theta'_i}$$

由于 $b(\cdot)$ 是均衡,所以, $z = b(\theta'_i) = b(\theta_i)$ 是这一条件的解,因此:

$$(n-1)\theta'_i f(\theta'_i)^{n-2} = \frac{\partial t^+(z)}{\partial z} \Big|_{z=\theta'_i}$$

这一微分方程的解为:

$$t^+(\theta'_i) = t^+(0) + \int_0^{\theta'_i} (n-1)\theta'_i f(\theta'_i)^{n-2} d\theta'_i$$

于是,在 $t^+(0) = 0$ 的假设下,我们有:

$$t^+(\theta'_i) = \int_0^{\theta'_i} (n-1)\theta'_i f(\theta'_i)^{n-2} d\theta'_i$$

它与拍卖形式无关。由于:

$$ER = \sum_{i=1}^n \int_0^k t_i(b(\theta_i)) dF(\theta_i) = n \int t^+(\theta'_i) f(\theta'_i) d\theta'_i$$

所以,命题得证。 $\square$

11.5 应用:竞选和全支付拍卖*

全支付拍卖要求所有竞拍者,无论是否赢得商品,都得支付自己的报价。许多竞赛具有这种特征。特殊利益集团行贿,其中行贿数量最多的团体得到自己想要的政策。候选人竞选公职,其中投入竞选经费最多的人胜出。这里探讨后一情况,并用收入等价原理分析不完善投票机制的效率特征。第6章分析了这一模型的一种特殊情况。

假设候选人集合为 $C = \{1, 2, \dots, c\}$ 。每个候选人都有一个非负的实值类型 $\theta_p \in \mathbb{R}_+$, 反映他在公职岗位上具有的效率。社会目标是选出最有效率的候选人 $p^* \in \arg\max_{p \in C} \theta_p$ 。质量类型 $\theta = (\theta_1, \dots, \theta_c)$ 是私人信息,只有候选人 p 才知道自己的类型 θ_p 。为简化分析,假设类型服从独立同分布,分布函数为 $F(\cdot)$, 它的密度函数为 $f(\cdot)$ 。 $M_c = \max\{\theta_j\}_{j=1}^c$ 表示 $c-1$ 次抽取 θ 所得到的最大类型。给定 θ_p , $M_c \leq \theta_p$ 的概率是 $F_c(\theta_p) \equiv F(\theta_p)^{c-1}$ 。求导,得到 M_c 的密度,为 $f_c(\theta_p) = (c-1)F(\theta_p)^{c-2}f(\theta_p)$ 。

我们的基本模型包括两个阶段。在第一阶段,每个候选人同时选择竞选努力程度 a_p 。在第二阶段,努力程度最高的候选人胜出。这里的关键假设是,由候选人的竞选努力程度导致的成本随其任职效率递减。就是说,如果候选人1的竞选成本低于候选人2,候选人1可能是更有效率的领导人。这一假设的合理性是,能力能够转化为在公职岗位上的卓越管理和竞选活动上的卓越纪录。为简化分析,我们假设竞选成本与效率成反比。候选人的收益于是为:

$$Eu(a_p, \theta_p) = \begin{cases} 1 - \frac{a_p}{\theta_p}, & \text{如果 } a_p > \max_{j \neq p} \{a_j\} \\ \frac{a_p}{\theta_p}, & \text{如果 } a_p < \max_{j \neq p} \{a_j\} \\ \frac{1}{\#\{j: a_j = a_p\}} - \frac{a_p}{\theta_p}, & \text{如果 } a_p = \max_{j \neq p} \{a_j\} \end{cases}$$

其中最后一行表达式说的是,如果出现平局,则随机选择最终胜出者。用 θ_p 乘以每个候选人的效用函数,上述收益函数就转化为了标准的全支付拍卖中的收益函数:竞拍物的价值为 θ_p ,参与竞拍要发生一个单位的成本。我们现在探讨对称的纯策略均衡,其中,每个候选人的努力程度 $\alpha(\theta_i)$ 是其效率类型 θ_i 的严格递增的可微函数。于是,类型为 θ_p 的候选人 p ,如果所选择的努力程度为 a_p ,得到的期望效用为:

$$\pi(a_p, \theta_p) = F(\alpha^{-1}(a_p))^{c-1} - \frac{a_p}{\theta_p}$$

最优的 a_p 必须是下面的一阶条件的解:

$$(c-1)F(\alpha^{-1}(a_p))^{c-2}f(\alpha^{-1}(a_p)) \frac{d\alpha^{-1}}{da_p} = \frac{1}{\theta_p}$$

在对称均衡中, $\alpha(\theta_p) = a_p$, 于是, $\alpha^{-1}(a_p) = \theta_p$ 。所以, 一阶条件可简化为:

$$\frac{d\alpha^{-1}}{da_p} = \frac{1}{\theta_p(c-1)F(\theta_p)^{c-2}f(\theta_p)}$$

或者,

$$\frac{d\alpha}{d\theta} = \theta_p(c-1)F(\theta_p)^{c-2}f(\theta_p)$$

积分, 得到了我们想要的解:

$$\begin{aligned}\alpha(\theta) &= \int_0^\theta x(c-1)F(x)^{c-2}f(x)dx \\ &= \int_0^\theta xf_c(x)dx\end{aligned}$$

这个函数在 θ 上严格递增。现在还需验证这个解满足充分二阶条件。这一结论得自 Krishna 和 Morgan(1997)中的定理 2, 这里不给出证明。

命题 11.1 存在对称均衡, 在均衡中, 努力程度 $\alpha(\theta_i)$ 在效率上严格递增。在这一均衡中, 最优秀的候选人 $p^* = \operatorname{argmax}_{p \in C} \theta_p$ 以概率 1 胜出。

有几点值得指出。第一, 由于 $\int_0^\theta xf_c(x)dx < \theta$, 所以, $\alpha(\theta_i) < \theta_i$, 就是说, 胜出的候选人得到严格正收益。第二, 努力函数对 c 的导数为:

$$\frac{\partial \alpha(\theta)}{\partial c} = \int_0^\theta x(c-1)\ln(F(x))F(x)^{c-2}f(x)dx$$

由于 $F(x) \leq 1$, 这个导数是负的。因此, 对 $c < c'$, $\alpha_c(\theta) > \alpha_{c+1}(\theta)$, 其中, $\alpha_n(\theta)$ 为 $c = n$ 时的均衡努力函数。这就是说, 每个候选人投入到竞选活动中的努力程度, 随候选人数量的上升而下降。第三, 这一模型等价于价值认定相互独立且 $\alpha(0) = 0$ 时的最高价格全支付拍卖, 所以, 根据收入等价原理, 期望总努力程度 $A_c = c \int_0^\infty \alpha(x)f(x)dx$, 肯定和价值认定服从分布 $F(\cdot)$ 的次高价格拍卖中赢家的期望支付相等。这个期望值, 就是取自 $F(\cdot)$ 的 c 次高值的期望, 即:

$$\text{总努力程度} = \int_0^\infty xc(c-1)(1-F(x))F(x)^{c-2}f(x)dx$$

这个值在 c 上递增。因此, 总努力程度在候选人数量上递增。

有理由认为, 竞选活动的社会目标是最大化 θ_{p^*} 。但是, 如果竞选活动的实际投入是在浪费社会资源, 则在提高胜出候选人的质量的期望值 $E\theta_{p^*}$, 与降低期望竞选成本 $E \sum_{p=1}^c \frac{\alpha(\theta_p)}{\theta_p}$ 之间, 必须权衡取舍。设 $v(\cdot)$ 和 $u(\cdot)$ 为二阶可微增函数, 且 $v'' < 0$, $u'' > 0$ 。我们分析下面的社会福利函数:

$$V_c = Ev(\theta_{p^*}) - Eu\left(\sum_{p=1}^c \frac{\alpha(\theta_p)}{\theta_p}\right)$$

竞选活动的全支付拍卖特征在选出最高质量的候选人方面,做得相当出色,但是,其中的成本 $\sum_{p=1}^c \frac{\alpha(\theta_p)}{\theta_p}$ 却未必合理。这就产生了一个问题:如何既能选出好的候选人又能防止竞选活动中的过度浪费呢?这一问题的答案涉及不完善选举——美国佛罗里达州,可能还有其他地方,似乎有时发生这种选举活动。

我们的基本模型假定选举是个完善的甄别机制,能选出最高质量的候选人。但是,在现实世界中,投票机制并不完善。在这里,我们探讨有随机结果的选举——有最高质量 a 的候选人以概率 $q + \frac{1-q}{c}$ 当选,其他候选人以概率 $\frac{1-q}{c}$ 当选。在这一背景中,候选人的收益为:

$$Eu(a_p, \theta_p) = \begin{cases} q + \frac{1-q}{c} - \frac{a_p}{\theta_p}, & \text{如果 } a_p > \max_{j \neq p} \{a_j\} \\ \frac{1-q}{c} - \frac{a_p}{\theta_p}, & \text{如果 } a_p < \max_{j \neq p} \{a_j\} \\ \frac{q}{\#\{j: a_j = a_p\}} + \frac{1-q}{c} - \frac{a_p}{\theta_p}, & \text{如果 } a_p = \max_{j \neq p} \{a_j\} \end{cases}$$

给定这一效用函数,均衡策略的推导过程类似基本模型中的推导过程(在基本模型中, $q = 1$)。如果每个候选人使用严格递增的可微策略 $\alpha(\theta_i)$,则类型为 θ_p 的候选人 p ,从选择 a_p 中得到的期望效用为:

$$\pi(a_p, \theta_p, q) = qF(\alpha^{-1}(a_p))^{c-1} - \frac{a_p}{\theta_p} + \frac{1-q}{c}$$

最优的 a_p 必须是下面的一阶条件的解:

$$q(c-1)F(\alpha^{-1}(a_p))^{c-2}f(\alpha^{-1}(a_p)) \frac{d\alpha^{-1}}{da_p} = \frac{1}{\theta_p} \quad (11.2)$$

对所有的 q , 对称均衡映射要求 $\alpha(\theta_p; q) = a_p$, 于是, $\alpha^{-1}(a_p; q) = \theta_p$ 。(11.2)式因此可简化为:

$$\frac{d\alpha^{-1}}{da_p} = \frac{1}{\theta_p q(c-1)F(\theta_p)^{c-2}f(\theta_p)}$$

或者说,

$$\frac{d\alpha}{d\theta} = \theta_p q(c-1)F(\theta_p)^{c-2}f(\theta_p)$$

积分,得到我们所求的解:

$$\begin{aligned} \alpha(\theta; q) &= \int_0^\theta x q(c-1)F(x)^{c-2}f(x)dx \\ &= q \int_0^\theta x f_c(x)dx \\ &= q\alpha(\theta; 1) \end{aligned} \quad (11.3)$$

因此,在选举不完善时,每个候选人都会根据不完善性参数 q ,成比例地降低自己的竞选努力程度。

我们现在探讨只有两个候选人参与竞选的情况,并求解最优的 $q \in [0, 1]$ 。在这种情况下,两个人的策略是:

$$\alpha(\theta; q) = q \int_0^\theta x f(x) dx$$

与 q 相对应的社会福利为:

$$\begin{aligned} V(q) &= q \int_0^\infty v(x) 2F(x) f(x) dx \\ &\quad + (1-q) \int_0^\infty v(x) 2(1-F(x)) f(x) dx \\ &\quad - \int_0^\infty u\left(\frac{2q}{\theta} \int_0^\theta x f(x) dx\right) df(\theta) \end{aligned}$$

q 的最优值满足一阶条件:

$$\begin{aligned} &\int_0^\infty v(x) F(x) f(x) dx - \int_0^\infty v(x) (1-F(x)) f(x) dx \\ &= \int_0^\infty \left(\left[\frac{1}{\theta} \int_0^\theta x f(x) dx \right] u' \left(\frac{2q}{\theta} \int_0^\theta x f(x) dx \right) \right) f(\theta) d\theta \end{aligned}$$

如果:

$$\begin{aligned} &\int_0^\infty v(x) F(x) f(x) dx - \int_0^\infty v(x) (1-F(x)) f(x) dx \\ &< \int_0^\infty \left(\left[\frac{1}{\theta} \int_0^\theta x f(x) dx \right] u' \left(\frac{2}{\theta} \int_0^\theta x f(x) dx \right) \right) f(\theta) d\theta \end{aligned}$$

最优值 q^* 小于 1;

如果:

$$\begin{aligned} &\int_0^\infty v(x) F(x) f(x) dx - \int_0^\infty v(x) (1-F(x)) f(x) dx \\ &> \int_0^\infty \left(\left[\frac{2}{\theta} \int_0^\theta x f(x) dx \right] u'(0) \right) f(\theta) d\theta \end{aligned}$$

最优值 q^* 大于 0。

就是说,如果 $v(\cdot)$ 相对平坦而 $u(\cdot)$ 相对陡峭,则不完善选举能提高效率。如果 v 和 u 都为线性,则一阶条件不依赖 q ,最优的 q 或者是 1 或者是 0——有效率的做法或者是最大化最好的候选人当选的概率,或者是最小化竞选成本。

11.6 激励一致性和个人理性

在购买咖啡机例子和民调例子所探讨的机制中,诚实地报告自己的类型构成

最优反应。揭示原理指出,只探讨直接机制不会减少所面对的选择函数的数量,但是,它没有指出哪些类型的选择函数是可执行的。换句话说,它没有回答我们的问题:哪些类型的机制诱使代理人说真话?激励一致性条件要求在给定的机制下,代理人*i*如果推断其他人都说真话,也会选择说真话而不是说谎。

设代理人集合为*N*,选择空间为*X*,类型空间为 Θ ,类型空间上的先验联合密度函数为 $f(\theta)$,且对每个*i* ∈ *N*,状态依赖效用函数为 $u_i(x, \theta)$ 。直接机制为映射 $p(\theta): \Theta \rightarrow X$ 。在由这一机制诱生的博弈中,要有使用说真话策略的贝叶斯 Nash 均衡,以下激励一致性条件必须得到满足:

定义 11.4 (激励一致性)对每个*i* ∈ *N* 和 Θ_i 中每对不同的 θ_i 和 θ'_i ,

$$\int_{\Theta_{-i}} u_i(p(\theta_i, \theta_{-i}), \theta_i) f_{-i}(\theta_{-i}) d\theta_{-i} \geq \int_{\Theta_{-i}} u_i(p(\theta'_i, \theta_{-i}), \theta_i) f_{-i}(\theta_{-i}) d\theta_{-i} \quad (\text{IC})$$

这个激励一致性条件(或者说,条件 IC)要求,在其他代理人发出说真话的信号时,代理人*i* ∈ *N* 发出说真话的信号构成最优反应。保证这个 IC 条件得到满足的最简单的方式,是找到一个机制,其中, $\int_{\Theta_{-i}} u_i(p(m_i, \theta_{-i}), \theta_i) f_{-i}(\theta_{-i}) d\theta_{-i}$ 在 m_i 上为常数。我们看下面的例子。设 ω 是个随机状态变量,它影响参与者从各种政策中得到的收益。假设,有 $n \geq 3$ 个代理人观察到 ω ,就是说,以 1 的概率, $\theta_i = \omega$ 。机制设计者可以利用下面的简单机制,执行选择函数 $x(\omega)$ ——它把信号组合映射到政策 x 上:

$$p(\theta) = \begin{cases} x(\omega'), & \text{如果 } \#\{j \in N: \theta_j = \omega'\} \geq n-1 \\ x(\omega^*), & \text{在其他情况中} \end{cases}$$

其中, $\#\{j \in N: \theta_j = \omega'\}$ 表示宣布 $\theta_i = \omega'$ 的代理人的数量, ω^* 是个任意的 ω 值。由于一次偏离并不改变这一政策选择,所以,这一机制满足激励一致性。

但是,这一机制没有为说真话提供强大激励。在拍卖理论中,激励一致性在次高价格拍卖中得到满足,因为期望效用(这里不考虑出现平局的可能性)为:

$$u_i(p(m_i, \theta_{-i}), \theta_i) = \begin{cases} \theta_i - \max_j \theta_j, & \text{如果 } m_i > \max_j \theta_j \\ 0, & \text{如果 } m_i \leq \max_j \theta_j \end{cases}$$

在 $m_i > \max_j \theta_j$ 时,这一函数在 m_i 上是常数;在 $m_i < \max_j \theta_j$ 时,这一函数在 m_i 上也是常数。这一函数中可变的从 0 跳到 $\theta_i - \max_j \theta_j$, 而只有当 $m_i > \max_j \theta_j$ 时, $\theta_i - \max_j \theta_j$ 才是正数。与此不同的是,最高价格拍卖不满足激励一致性。换句话说,在最高价格拍卖中,说真话的报价策略不构成贝叶斯 Nash 均衡策略。

我们以购买咖啡机问题为例,直观地讨论激励一致性条件。设 $t_i(m_i, m_{-i})$ 是任意支付函数,它把信号组合映射到第*i* 个人需支付的金额上。设 $p(m)$ 是政策函数,它把信号组合映射到购买咖啡机的概率上。给定这一支付函数,第*i* 个教师从宣布 m_i 的行为中得到的期望效用为:

$$Eu(m_i, \theta_i) \equiv \int [\theta_i p(m_i, \theta_{-i}) - t_i(m_i, \theta_{-i})] f_{-i}(\theta_{-i}) d\theta_{-i}$$

根据激励一致性条件, $m_i = \theta_i$ 最大化 $Eu(m_i, \theta_i)$ 。如果 p 和 t 是可微函数, 我们就能以一阶条件为基础, 分析局部激励一致性条件。一阶条件为:

$$\left. \frac{\partial Eu(m_i, \theta_i)}{\partial m_i} \right|_{m_i = \theta_i} = 0$$

交换积分和微分顺序, 得到:

$$\theta_i \int \frac{\partial p(m_i, \theta_{-i})}{\partial m_i} f_{-i}(\theta_{-i}) d\theta_{-i} \Big|_{m_i = \theta_i} = \int \frac{\partial t(m_i, \theta_{-i})}{\partial m_i} f_{-i}(\theta_{-i}) d\theta_{-i} \Big|_{m_i = \theta_i}$$

直观地看, 这一条件要求在激励一致性机制中, 略为低报 θ_i 所导致的支付的预期减少, 正好被购买咖啡机的预期概率的下降与价值认定 θ_i 的乘积抵消。

激励一致性条件要求参与者愿意揭示其私人信息, 但是, 要执行选择函数, 参与者还必须愿意进行这一博弈。对这一类约束条件的分析有时更简单, 也更具体。关键的假设是, 某个参与者, 只有在他从均衡策略中得到的期望收益至少等于他不参与博弈所能得到的期望收益时, 才会理性地参与博弈。在一些问题中, 这类约束条件并不重要: 代理人别无选择而只能参与博弈。在其他问题中, 我们可能需要找到在代理人不参与博弈时, 他所实现的价值。正式地说, 我们有以下约束条件。

定义 11.5 (个人理性约束) 对每个 $i \in N$, 设不参与某个机制时所能实现的价值为 $v_i(\theta_i, \theta_{-i})$, 在直接机制中说真话所能实现的价值为 $u_i(\theta_i, \theta_{-i})$ 。对每个 $i \in N$ 和每个 $\theta_i \in \Theta$,

$$\int u_i(\theta_i, \theta_{-i}) dF(\theta_{-i}) \geq \int v_i(\theta_i, \theta_{-i}) dF(\theta_{-i}) \quad (\text{IR})$$

还有个更弱的个人理性(或 IR)约束条件, 它只要求参与者在事前(在他了解到自己的类型前)愿意进行这一博弈。它被称为事前个人理性约束, 而前面所定义的条件被称为期中个人理性约束。^①

11.7 约束条件下的机制设计

在标准的机制设计问题中, 可供设计者使用的机制数量很多。激励一致性和个人理性是唯一的约束条件。但是, 在许多政治背景中, 委托人常常无法使用某些机制或者无法使用某些激励手段如直接的货币奖励等。在本章剩余部分中, 我们通过一些例子, 介绍政治科学家是如何利用约束条件下的机制设计, 解决制度设计问题的。

政治科学中的模型一般落于委托代理理论范畴内: 有委托人或老板, 还有代理人或下级。委托人喜欢“好的”代理人, 希望他们“做好工作”。但是, 监督问题使委

^① 在这个概念和重复博弈分析中所使用的个人理性概念之间, 可能存在混淆。文献同时以这两种相互有联系的但是又彼此不同的方式使用这个术语。

托人无法直接观察到代理人是否是“好的”或者是否在“做好工作”。另一方面,代理人想使委托人相信他们是“好的”,在“做好工作”。但是,代理人一般偏好偷懒而不是做好工作。在大部分政治科学模型中,委托人有一些手段可影响代理人的行为。用机制设计的术语来说,委托人是设计者,代理人“做好工作”或者揭示自己是否是“好的”类型,表现为在直接机制中选择恰当的信号。委托人可利用的手段表现为可实行的机制上的约束条件。在探讨更有意义的应用之前,我们用一个简单的委托模型说明这些概念。

假设有一个委托人和两个代理人。代理人的类型为 $\theta_i \in \{\text{好的, 差的}\}$ 。每个代理人为“好的”类型的概率是 π 。一个代理人,在进行某一任务时,所投入的努力程度为 $a_i \in \{\text{高, 低}\}$ 。如果为 $a_i = \text{高}$,则代理人要发生成本 c ;如果 $a_i = \text{低}$,则成本等于 0。委托人必须从两人中选出一位完成一项任务,被选中的代理人必须决定自己选择哪一努力水平。委托代理模型(或者说,代理模型)的根本问题可以表述为:委托人如何通过对制度的设计,选出好的代理人,并诱使他选择高的努力程度呢?其中,通过对制度的设计选出好的类型,所解决的是逆选择问题;通过对制度的设计激励代理人投入高的努力水平,所解决的是道德风险问题。

如果委托人能够观察到代理人的类型,譬如说,代理人的衬衫上贴有标签,标明了他的类型,如果代理人对努力程度的选择能够被观察到,就不存在监督问题或者可观察性问题。在这种情况下,委托人的工作简单之至。他只需雇用好的类型,并在他们不肯卖力工作时,解雇他们。但是,一般情况下,在代理人的类型上存在着私人信息,并且代理人的努力程度无法被完善地观察到。

11.7.1 选举

Ferejohn(1986)把选举作为委托代理问题进行分析。我们看 Ferejohn 模型的一个简化版本。Ferejohn 只探讨了道德风险问题,所以,我们暂时不考虑代理人类型上的差异。假设两个政党完全相同。执政党选择了某一努力程度,但是,选民只观察到包含噪音的政策结果 $x \in \{\text{高, 低}\}$ 。如果 $a = \text{高}$,则 $x = \text{高}$ 的概率为 $q > \frac{1}{2}$, $x = \text{低}$ 的概率为 $1 - q$; 如果 $a = \text{低}$,则 $x = \text{低}$ 的概率为 q , $x = \text{高}$ 的概率为 $1 - q$ 。最后,假设执政党在每个时期中得到租金 r ,并以贴现因子 δ 贴现未来。在野党每个时期中得到的收益为 0。

选民选择选举规则,规则中规定使执政党自然连任的 x 值。那么,选民是否能设计一种规则,激励执政党始终选择 $a = \text{高}$ 吗?一种简单的规则是,如果 $x = \text{高}$,就让执政党连任,否则就撤换它。如果在每个时期里,选民都使用这一规则,执政党就面临着一个简单的决策问题。如果它在当前时期中选择 $a = \text{高}$,则得到的期望效用为:

$$r - c + q\delta V_1 + (1 - q)\delta V_0 \quad (11.4)$$

在下个时期开始时,如果当前的执政党依然在位,则自下一个时期开始的博弈带给它的价值为 V_I ;在下个时期开始时,如果它不再执政,则自下一个时期开始的博弈带给它的价值收益为 V_O 。^①如果执政党在当前时期中选择 $a = \text{低}$,则它的期望效用为:

$$r + (1-q)\delta V_I + q\delta V_O \quad (11.5)$$

如果选民使用上述规则,执政党会发现,在每个时期中都选择 $a = \text{高}$ 是最优做法,这一博弈带给它的价值于是为:

$$V_I = r - c + q\delta V_I + (1-q)\delta V_O$$

$$V_O = (1-q)\delta V_I + q\delta V_O$$

求解这四个方程,得到:

$$V_I = \frac{r - c + q\delta V_I + (1-q)\delta V_O}{2q\delta^2 - \delta^2 - 2q\delta + 1} \quad (11.6)$$

$$V_O = \frac{r\delta - c\delta + q\delta V_I + q\delta V_O}{2q\delta^2 - \delta^2 - 2q\delta + 1} \quad (11.7)$$

对执政党来说,选择 $a = \text{高}$ 的激励一致性条件要求(11.4)式不小于(11.5)式。因此,这一激励约束为:

$$V_I - V_O \geq \frac{c}{(2q-1)\delta}$$

将(11.6)式和(11.7)式代入到上式中,这一条件变为:

$$r - c - r\delta + c\delta \geq \frac{c(2q\delta^2 - \delta^2 - 2q\delta + 1)}{(2q-1)\delta} \quad (11.8)$$

投票规则要在每个时期中诱生出 $a_i = \text{高}$,条件(11.8)就必不可少。如果外生参数满足(11.8)式,则只在 $x = \text{高}$ 时才让执政党连任的规则,能够解决道德风险问题。由于选民在 $x = \text{高}$ 和 $x = \text{低}$ 之间,偏好前者,所以,选民会坚持这一规则:给定执政党的决策规则,这一规则构成均衡策略。由于所有执政党都选择 $a_i = \text{高}$ 并且所有执政党都有相同类型(就是说,不存在逆选择问题),所以这一结论肯定成立。因此,如果上述条件得到满足,我们就得到了两个政党和选民之间的这个博弈的 Nash 均衡:执政党选择 $a_i = \text{高}$,只在 $x = \text{高}$ 时,选民才让执政党连任。

我们在逆选择背景下分析这一问题。假设执政党并不选择自己的努力水平,而是拥有关于自己类型的私人信息。为分析简单,假设每个政党的类型为“高”的概率是 $\pi > \frac{1}{2}$,为“低”的概率是 $1 - \pi$ 。两个政党的类型是独立的。在每个时期中,

① 见第3章中对 Bellman 方程和值函数的讨论。

选民只观察到 x 。如果执政党的类型为“高”，则以概率 $z > \frac{1}{2}$ ， $x = \text{高}$ ；以概率 $1 - z$ ， $x = \text{低}$ 。如果执政党的类型为“低”，则以概率 $1 - z$ ， $x = \text{高}$ ；以概率 z ，“ $x = \text{低}$ ”。选民依然可以使用前述简单规则：如果 $x = \text{高}$ ，就让执政党连任；否则，就撤换。但是，这一规则可能会让暂时不走运的高质量政府下台。更好的规则应利用以前时期的信息帮助做出判断。从党派 i 执政的时期中，任选出设 k_i 个有限时期， h_i 为在这些时期中 $x = \text{高}$ 的时期的数量。于是，在这 k_i 个时期中，有 $k_i - h_i$ 个时期， $x = \text{低}$ 。在 k_i 个时期中，如果有 h_i 个时期里 $x = \text{高}$ ，则根据贝叶斯法则，党派 i 的类型为“高”的后验概率为：

$$\Pr(\theta_i = \text{高} \mid k_i, h_i) = \frac{\pi z^{h_i} (1-z)^{k_i-h_i}}{\pi z^{h_i} (1-z)^{k_i-h_i} + (1-\pi) (1-z)^{h_i} z^{k_i-h_i}}$$

最优选举规则是，让执政党 i 继续执政，直至 $\Pr(\theta_i = \text{高} \mid k_i, h_i)$ 小于党派 j 的事前预期质量 π 。此时，让政党 j 上台。接着，只要 $\Pr(\theta_i = \text{高} \mid k_i, h_i) > \Pr(\theta_j = \text{高} \mid k_j, h_j)$ ，就让党派 j 下台。这里对代理理论的介绍，只是起个导论的作用。我们现在讨论代理理论在分析授权中的几个应用。

11.7.2 授权模型

在现代行政体制中，一个重要的问题是对行政机构的政治控制权与此机构为应用其专业才能解决政策问题所需的自主权之间的关系。委托代理理论是研究这一问题的自然不过的方法。它在政治学中的最早应用，是研究美国国会在什么条件下以什么方式把一些立法权下放给规制机构，自此之后，它就应用于对其他许多政治体制的研究中，并应用于对其他许多政治机构包括政党和国际组织的研究中。本节中所介绍的模型，曾被 Epstein 和 O'Halloran(1994)用于研究美国的立法权下放问题。

假设立法机构 L 在考虑把多少权力下放给行政机构 A 。政策空间 X 是 $\mathbb{R}$ 的子集。在 n 个立法议员中，每个人都有二次方政策偏好，理想点分别为 $l_1 < \dots < l_n$ 。行政机构 A 被看作是独立的行为主体，有自己的二次方政策偏好，理想点为 a 。

立法议员不知道各种政策 $p \in X$ 所导致的结果，而在这一点上， A 拥有充分信息。我们对这种不确定性的处理仿照 Gilligan 和 Krehbiel(1989)在对立法委员会的研究中所使用的方法。假设政策结果 x 是政策选择 p 和误差项 ε 的函数。为使分析尽可能简单，设 $x = p - \varepsilon$ 。每个立法议员事前都知道 ε 的均值为 0，并服从分布函数 $F(\varepsilon)$ ，它的密度为 $f(\varepsilon)$ 。行政机构 A 则确定地知道 ε 。在这种信息结构中，行政机构是政策专家，立法议员依靠它提高政策制定效率。

我们不从 L 的角度求解最优机制，而是分析一类简单的机制： L 选择可下放给行政机构的政策集 $P \subset X$ 。依照文献中的标准处理方法，立法机构所选择的 P 在多数通过的规则下不会被投票推翻。行政机构如果越权选择 $p \notin P$ ，就会受到惩罚。

惩罚力度之强会使得这种越权行为永远不会发生。给定可下放的政策集 P 和自然状态 ε , A 通过对 p 的选择求解如下问题:

$$\max_p \{-(a-p+\varepsilon)^2\} \quad \text{s. t.} \quad p \in P$$

就是说,只要 $a+\varepsilon \in P$, 行政机构 A 就能通过选择 $p = a+\varepsilon$ 得到自己的理想点。如果 $a+\varepsilon \notin P$, A 就从 P 中选择最靠近 $a+\varepsilon$ 的点。

给定 A 的最优反应,我们分析立法机构对 P 的选择。理论上说, P 可以是 X 的任意子集。但是,最优的 P 始终是个闭区间 $[\underline{p}, \bar{p}] \subset X$ 。

命题 11.2 P^* 是闭区间 $[\underline{p}, \bar{p}] \subset X$ 。

对这一命题的详细证明,见 Gailmard(2005)。这里只介绍证明思路。假设在 P 中有个“窟窿”(例如, $P = [\underline{p}, p'] \cup [p'', \bar{p}]$, 其中, $p' < p''$)。则只要 $p' < a + \varepsilon < \frac{p' + p''}{2}$, A 就选择 p' ; 只要 $\frac{p' + p''}{2} < a + \varepsilon < p''$, A 就选择 p'' 。就是说,政策结果作为 ε 的函数,表现为图 11.1 中的实线。我们不难看出,让 p' 和 p'' 接近,政策结果上的方差缩小。由于所有议员都厌恶风险,所以,只要预期政策结果没有变化,他们就会降低这一方差。你应该能够证明,在不改变平均政策结果的前提下,总有办法使 p' 和 p'' 彼此接近。^①

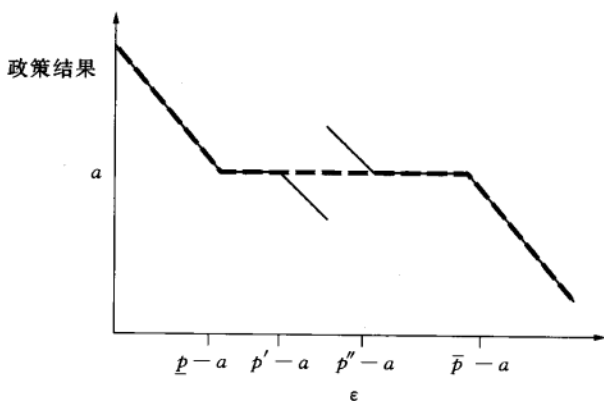


图 11.1 Epstein-O'Halloran 模型中的政策结果

根据这一结论,立法机构的集体选择问题显然是选择 $\underline{p}$ 和 $\bar{p}$ 。行政机构的最优反应函数因此为:

$$p^* = \begin{cases} \underline{p}, & \text{如果 } a + \varepsilon < \underline{p} \\ \bar{p}, & \text{如果 } a + \varepsilon > \bar{p} \\ a + \varepsilon, & \text{在其他情况中} \end{cases}$$

① 去掉区间中的“窟窿”并不意味着这个区间肯定是闭的。这个区间要为闭集,要求 A 的最优反应有明确界定。

图 11.1 中的虚线就是行政机构的最优反应函数。

对 L 的集体选择的分析,则有点儿复杂,因为它必须在两个方面做出选择,既要选择区间下界,又要选择区间上界。幸运的是,多数通过的规则下产生的决策,是有中位理想点 l_m 的议员所偏好的 P^* 。我们来证明这一判断。给定 A 的最优反应, l_i 所偏好的 $\underline{p}$ 与 $\bar{p}$ 最大化以下目标函数:

$$-\int_{\bar{p}-a}^{\infty} (l_i - \bar{p} + \varepsilon)^2 f(\varepsilon) d\varepsilon - \int_{\underline{p}-a}^{\bar{p}-a} (l_i - a)^2 f(\varepsilon) d\varepsilon - \int_{-\infty}^{\underline{p}-a} (l_i - \underline{p} + \varepsilon)^2 f(\varepsilon) d\varepsilon$$

上式分别对 $\bar{p}$ 和 $\underline{p}$ 求导^①,得到两个一阶条件:

$$2 \int_{\bar{p}-a}^{\infty} (l_i - \bar{p} + \varepsilon) f(\varepsilon) d\varepsilon = 0$$

$$2 \int_{-\infty}^{\underline{p}-a} (l_i - \underline{p} + \varepsilon) f(\varepsilon) d\varepsilon = 0$$

只要 $l_i < a$, 则对任何有限的 $\underline{p}$, $\frac{\partial}{\partial \underline{p}} < 0$ 。因此,最优选择是 $\underline{p}^* = -\infty$ 。为什么呢? 当 $l_i < a$ 时,行政机构想要的政策,比议员 i 在拥有充分信息时想要的政策要高。就是说,立法议员 i 会发现,禁止 A 选择低政策永远不符合自己的利益。如果 $l_i > a$, 则 $\frac{\partial}{\partial \bar{p}} > 0$, 于是 $\bar{p}^* = \infty$ 。综合这两个结论,我们就能看到,任何政策,只要对有理想点 l_m 的议员不是最优的,在多数通过的投票机制下就无法得到通过。

命题 11.3 多数通过规则下的结果 P^* , 是有理想点 l_m 的立法议员所偏好的闭区间 $[\underline{p}, \bar{p}]$ 。

由于多数通过的投票规则产生的结果是中位议员的理想点,所以,权力下放博弈变为了行政机构和立法议员之间的博弈。于是,可授权给行政机构的政策集,在 $l_m < a$ 时,由以下两个条件决定:

$$\underline{p}^* = -\infty$$

$$\int_{\bar{p}^*-a}^{\infty} (l_m - \bar{p}^* + \varepsilon) f(\varepsilon) d\varepsilon = 0$$

在 $l_m > a$ 时,由以下两个条件决定:

$$\bar{p}^* = \infty$$

$$\int_{-\infty}^{\underline{p}^*-a} (l_m - \underline{p}^* + \varepsilon) f(\varepsilon) d\varepsilon = 0$$

① 这里的求导用到了莱布尼茨法则。

我们分析其中的有限条件。如果 $l_m < a$, 我们可以把 $\bar{p}^*$ 的表达式写为

$$l_m = \int_{\bar{p}^* - a}^{\infty} (\bar{p}^* - \varepsilon) \frac{f(\varepsilon)}{1 - F(\bar{p}^* - a)} d\varepsilon$$

这一条件指出, 在施加于 A 上的约束的限制下, 期望结果必须等于中位议员的理想点。如果这一期望结果大于 l_m , 中位议员能通过进一步限制 A 的选择, 得到更低的政策, 从而在期望意义上得到更好的结果。我们可以在 $l_m > a$ 时把 $\underline{p}^*$ 上的条件写为:

$$l_m = \int_{-\infty}^{\underline{p}^* - a} (\underline{p}^* - \varepsilon) \frac{f(\varepsilon)}{1 - F(\underline{p}^* - a)} d\varepsilon$$

为得到更具体的结论, 设 ε 在 $[-E, E]$ 上均匀分布, 于是 $F(\varepsilon) = \frac{\varepsilon + E}{2E}$,

$f(\varepsilon) = \frac{1}{2E}$ 。于是, 在 $l_m < a$ 时, 我们能得到 $\bar{p}^* = 2l_m - a + E$ 。从这一表达式中可以看出, 如果行政机构和立法机构通过提高 l_m 或者降低 a 而更加接近, 则立法机构通过的授权范围更加宽容, 就是说, 有更大的上界。形象地说, 立法机构会下放更多的权力给与它有相同偏好的行政机构。当政策上的不确定性更严重时 (例如, E 更大时), 行政机构也能得到更大权力。也就是说, 在信息不对称性更严重时, 立法机构更依赖拥有信息的行政机构制定政策。^①

11.7.3 行政能力

在 Epstein 和 O'Halloran (1994) 模型中和在其他的授权模型中, 一个重要的假设是, 行政机构能够无偏差地完善地执行它的政策选择。对发达国家来说, 这可能是一个合理假设, 因为它拥有受过严格专业训练的行政官员; 但是, 对发展中国家或已经逝去的历史时代而言, 这一假设可能并不适用。

由于标准模型的这种局限性, Huber 和 McCarty (2004) 提出了新的模型, 假定行政官员们执行政策的能力参差不齐。在这个模型中, A 试图执行政策 p , 但是所实现的政策是 $\tilde{p} = p - \omega$, 其中 ω 是政策执行误差, 它的均值为 0, 方差为 σ_ω^2 。能力强的官员能更好地执行政策, 因此其 σ_ω^2 较低。能力弱的官员在执行政策时准确性较低, 因此 σ_ω^2 较高。设 $G(\omega)$ 是 ω 的分布函数, $g(\omega)$ 是密度函数。

Huber 和 McCarty 把这个能力模型嵌入到非常类似 Epstein 和 O'Halloran 的授权模型中。立法机构想授权给行政机构, 因为后者对各种政策选择的结果拥有更好的信息。与前面一样, 行政机构知道 ε 而立法机构只知道 ε 在 $[-E, E]$ 上均匀分布。 L 的成员和 A 在政策空间 X 上都有二次方偏好。我们依然使用前一节中的符号。

① 如果 $l_m > a$, 解为 $\underline{p}^* = 2l_m - a - E$ 。你应该能够证明, 在 l_m 和 a 接近时, 或者在 E 更大时, 这一下界的限制性会放松。

立法机构首先行动,通过立法授权给行政机构的政策集为 $P = [\underline{p}, \bar{p}]$ 。^①接着,行政机构试图执行政策 p ,却导致了政策 $\tilde{p} = p - \omega$ 和结果 $x = \tilde{p} - \varepsilon$ 。根据立法,必须有 $\tilde{p} \in P$ 。这就意味着,即使行政机构想遵守立法授权规定(即选择 $p \in P$),执行误差依然可能导致不合规定的结果。如果 $\tilde{p} \notin P$,作为对违规结果的惩罚,行政机构要发生成本 δ 。^②与 Epstein 和 O'Halloran 不同,在这一模型中,行政机构可以有意违反规定。当然了,违规结果也可能由执行误差导致。由于我们假定 L 无法观察到 p ,所以,不管最终原因是什么,对违规结果的惩罚是一样的。就是说,在 $p - \omega > \bar{p}$ 时或者在 $p - \omega < \underline{p}$ 时, A 都要受到惩罚。给定对 p 的选择,惩罚概率为 $G(p - \bar{p}) + 1 - G(p - \underline{p})$ 。

为便利分析,我们把 Epstein 和 O'Halloran 模型的许多特征融入到这里的模型中。首先,我们能够证明,多数通过规则下所选择的 P 最大化有理想点 l_m 的立法议员的效用。这里始终假定 $a > l_m$ 。其次,在当前的模型中,如果 $a > l_m$,我们能够证明 $p^* = -\infty$ 。

给定这一背景和 $a > l_m$ 的假定, A 的效用函数为^③:

$$\begin{aligned} & - \int_{-\infty}^{\infty} (a - p + \omega + \varepsilon)^2 g(\omega) d\omega - \delta [G(p - \bar{p})] \\ & = - (a - p + \varepsilon)^2 - \sigma_{\omega}^2 - \delta [G(p - \bar{p})] \end{aligned}$$

这个函数取最大值的一阶条件为:

$$2(a - p + \varepsilon) - \delta g(p - \bar{p}) = 0$$

在这一表达式中,第一项为期望政策趋近行政机构的理想点 a 时所实现的边际收益;第二项为提高政策所导致的净边际成本,是用惩罚概率表示的。提高 p 使得 $p + \omega > \bar{p}$ 的概率上升 $g(p - \bar{p})$ 。 A 通过对 p 的选择使边际政策收益等于边际惩罚成本。为得到显性解 p^* , Huber 和 McCarty 假定 $g(\omega) = \frac{\Omega - |\omega|}{\Omega^2}$ 。这个密度函数在区间 $[-\Omega, \Omega]$ 上呈“伞状”,并满足条件 $\sigma_{\omega}^2 = \frac{\Omega^2}{6}$ 。就是说, Ω 是衡量行政机构无能的指标。

- ① 阅读原始论文的读者需要注意,在 Huber 和 McCarty 的论文中,政策委托人是 一个一般意义上的政治家而不是立法机构。我们调整了有关术语和符号以便与我们对 Epstein 和 O'Halloran 的讨论相一致。
- ② Huber 和 McCarty 假定,违规现象能以一定的概率被察觉到。但是,这里的简化不影响模型的结论。
- ③ 这一表达式中的第二行来自关于二次方函数的期望效用的一个有用的特征:

$$\int_{-\infty}^{\infty} (x + \phi)^2 f(\phi) d\phi = (x - E(\phi))^2 + \text{var}(\phi)$$

其中, E 是期望算子, var 是方差运算符。

给定关于函数形式的这些假设,图 11.2 给出了边际政策收益曲线,并对三个不同的 $\bar{p}$ 值,给出了边际服从成本曲线,它们都是 p 的函数。边际政策收益曲线是图中向下倾斜的粗直线。边际政策收益与 $\bar{p}$ 点的位置无关;随着行政机构的行动接近它最偏好的行动 $a + \epsilon$, 边际政策收益下降。边际成本曲线决定于 $\bar{p}$ 的位置,在图 11.2 中,是三个分别以 $\bar{p}_1$ 、 $\bar{p}_2$ 和 $\bar{p}_3$ 为中心的三角形。这些三角形表示函数 $\delta g(p - \bar{p})$ 。如果 $\bar{p}$ 太高或太低,则在行政机构的理想行动 $a + \epsilon$ 上,边际成本为 0。此时,行政机构选择的正是自己的理想行动。如果 $\bar{p}$ 值大小适中,则行政机构的最优反应位于边际收益曲线和相应的边际成本曲线的交点上。^①例如, $\bar{p}_1$ 导致最优行动 p_1^* 。对任意的 $p > p_1^*$, 提高政策行动(使政策趋向行政机构最偏好的政策)所实现的边际政策收益超过边际服从成本;对任意的 $p < p_1^*$, 政策远离行政机构最偏好的政策所导致的服从成本的下降,超过政策损失。

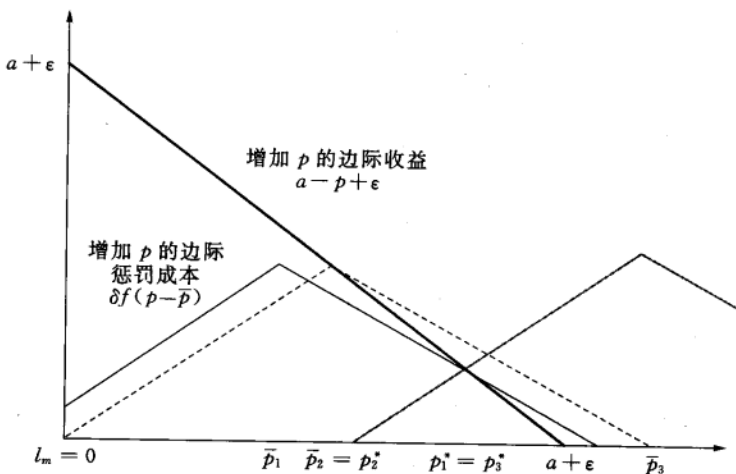


图 11.2 Huber-McCarty 模型中的政策选择

从图 11.2 中可以看出, $\bar{p}$ 上的变化对行政机构的最优反应的影响, 决定于“服从成本”三角形的顶点是位于政策收益线的左边还是右边(在图 11.2 中, 是位于 $\bar{p}_2$ 的左边还是右边)。在 $\bar{p}_1 < \bar{p} < \bar{p}_2$ 时, $p^* > \bar{p}$ 。在这一范围里, $\bar{p}$ 如果上升, 任意的 $p^* > \bar{p}$ 的边际服从成本都会上升, 诱使行政机构靠拢立法机构的理想点。但是, 在 $\bar{p}_2$ 上, 这一效应逆转。在 $\bar{p} > \bar{p}_2$ 时, $p^* < \bar{p}$ 。 $\bar{p}$ 如果上升, 任意的 $p^* < \bar{p}$ 的边际服从成本都会下降, 诱使行政机构向自己的理想点调整自己的行动。因此, 立法机构能够诱生的最小行动为 $p^*(\bar{p}_2)$ 。

Huber 和 McCartney 指出, 行政机构的最大化问题的正式解为:

① McCartney 和 Huber 假定 $\Omega^2 > \delta$, 从而保证边际收益曲线和边际成本曲线有唯一交点, 并且这个交点为全局最大值。由于在他们所探讨的情况中, 行政能力和惩罚违规行为的能力都很低, 所以, 这个假定有其合理性。

$$p^* = \begin{cases} a + \varepsilon, & \text{如果 } \bar{p} - \varepsilon \leq a - \Omega \\ \frac{\Omega^2(a + \varepsilon) - \delta(\bar{p} + \Omega)}{\Omega^2 - \delta}, & \text{如果 } a - \Omega \leq \bar{p} - \varepsilon \leq a - \frac{\delta}{\Omega} \\ \frac{\Omega^2(a + \varepsilon) - \delta(\bar{p} - \Omega)}{\Omega^2 + \delta}, & \text{如果 } \bar{p} - \varepsilon \leq a + \Omega \\ a + \varepsilon, & \text{如果 } \bar{p} - \varepsilon \geq a + \Omega \end{cases}$$

A 的最优反应函数有几个特征值得强调。首先,对于极端规定 ($\bar{p} \leq a - \Omega + \varepsilon$ 或 $\bar{p} \geq a + \Omega + \varepsilon$), A 的最优反应是试图执行自己的理想点。如果立法规定过于宽松或者过于苛刻,在 A 的理想政策上,边际服从成本为 0。其次, l_m 只能诱生区间 $\left[a - \frac{\delta}{\Omega} + \varepsilon, a + \varepsilon\right]$ 中的政策。它无法诱生更低的政策,因为在更严格的规定下, A 可能更经常地违规。这个最低政策在 Ω 上递增。因此,在行政机构能力较低时, l_m 控制 A 的能力下降,因为,能力低的行政机构,不管它选择什么政策,在更多时间里不会服从规定。而且,违规概率对行政机构的选择并不十分敏感。因此,行政机构会执行接近自己的理想点的政策,因为自己的这种行为招致的边际惩罚很小。这个模型发现了关于低行政能力的一个重要结论:立法机构更难通过立法控制低能力行政机构。

Huber 和 McCarty 想研究的是低能力体制中的行政问题,因此探讨的是在行政能力足够低时的最优立法。^①在这些假设下,最优立法为:

$$\bar{p}^* = a - \frac{\delta}{\Omega} + \frac{\delta E}{\Omega^2}$$

和 Epstein-O'Halloran 模型一样, Huber-McCarty 模型认为,在 E 更大时, L 下放更多的权力。但是,两者在偏好分歧上的结论却正好相反。在 Huber-McCarty 模型中, a 和 l_m 之间的距离越大,立法规定越宽容,因为面对着更苛刻的规定,能力低的行政机构更可能选择自己的理想点。如果 a 远离 l_m ,这种违规行为给 L 造成的成本会很高。因此, L 会给予能力极低的行政机构更多自主权,从而更强地激励它服从规定。在 Huber-McCarty 的模型中,还有一个结论与标准模型不同。标准模型一般认为,如果事后惩罚足够严厉,委托人就会下放更多权力。但是,在 Huber-McCarty 模型中,更高的 δ 与更严厉的立法规定相联系。严厉的惩罚甚至诱使能力低的行政机构服从规定。因此, L 不必仅仅为诱使行政机构服从规定而下放更多自主权。

11.7.4 一般授权模型

在关于授权问题的大部分博弈论分析中,都有一些程式化假设。首先,委托人

① “足够低的”能力,指的是 $\Omega > \min\left\{E, \frac{\delta}{a}, \sqrt{\delta}\right\}$ 。

和代理人都厌恶风险。实际上,在大部分应用研究中,他们都有二次方偏好。其次,政策空间和冲击都是单维度的。第三,大部分模型假定政策结果是政策选择和冲击的可加性函数。这些假设虽然能让我们构造出基本模型,但是,Bendor 和 Meirowitz(2004)指出,它们也使得我们很难把基本模型推广到其他政治环境中。具体地说,这些假设使我们无法确定在授权决策中哪些因素更重要。例如,是对代理人的选择更重要呢?还是监督和控制机制更重要呢?这里介绍 Bendor 和 Meirowitz 的基本观点。

假设有 1 个委托人和 n 个代理人。^①所有代理人有 $\mathbb{R}^d$ 中的理想点,委托人的理想点则为向量 $\mathbf{0}$ 。结果上的偏好表示为效用函数 $u_i(x) = h(-\|\mathbf{x} - \mathbf{y}_i\|)$, 其中, $h(\cdot)$ 是严格递增的连续函数, $\|\mathbf{z}\|$ 是欧几里得范数, $\mathbf{y}_i$ 是 i 在 $\mathbb{R}^d$ 中的理想向量。单维度政策空间上的二次方偏好是这一假设的特殊情况。与 Epstein-O'Halloran 和 Huber-McCarty 的模型一样, Bendor 和 Meirowitz 假定委托人拥有的信息少于代理人,但是,他们(1)允许使用任意函数形式;并且(2)对应不同的政策选择,有不同的不确定性。就是说,对任意政策 p , 结果 $x(p)$ 是由条件分布 $F(x|p)$ 产生的随机变量。这就意味着在某些政策结果上,不确定性可能更大。委托人只知道条件分布族,而拥有信息的代理人确切地知道从 p 到 x 的映射。Bendor 和 Meirowitz 还用到了他们所谓的完全冲击吸收(perfect shock absorption)条件。根据这个条件,拥有信息的代理人能够通过选择恰当的 p , 实现任何政策结果 x 。Epstein 和 O'Halloran 的假定,即 $x = p + \varepsilon$, 成为了这一条件的一种具体情况。由于政策执行上的冲击,在 Huber-McCarty 模型中,这种冲击吸收假设只在期望意义上才成立。

在 Huber-McCarty 模型中,行政官员缺少执行其目标政策的能力, Bendor 和 Meirowitz 则把代理人的能力差异定义为专业才能上的差异。以概率 q_i , 代理人 i 了解到随机冲击,通过选择 p , 实现任意的 x 。以概率 $1 - q_i$, 代理人 i 不掌握关于随机冲击的信息,此时,他所知道的并不比委托人多。

在他们的基本授权模型中,委托人决定是否授权。委托人如果不授权,就根据自己关于政策冲击的先验推断选择政策。她如果授权,就选择一位代理人,给他政策选择上的完全自主权。这位代理人选择政策 p , 博弈结束。我们现在介绍其中的一些关键结论。就目前而言,假定所有代理人的专业才能都很高,即对所有的 $i \in N$, $q_i = 1$ 。

命题 11.4 委托人愿意授权给其理想点落于以 $\mathbf{0}$ 为圆心以 ε 为半径的闭球 $B(\varepsilon, \mathbf{0})$ 中的代理人。 $B(\varepsilon, \mathbf{0})$ 被称为授权集。

$B(\varepsilon, \mathbf{0})$ 能相当直接地构造出来。如果代理人 i 控制着政策选择,代理人会选择产生结果 $x = \mathbf{y}_i$ 的政策。因此,只有在委托人偏好结果 $\mathbf{y}_i$ 而不是由委托人的最优政策选择生成的概率分布时,委托人才会授权给代理人 i 。只要委托人面对着不

^① Bendor 和 Meirowitz 探讨了有多个委托人时的情况。

确定性,委托人亲自执行政策所得到的效用 u'_0 ,就小于以概率 1 得到结果 $x=0$ 时的效用 $h(0)$ 。委托人显然会授权给理想点为 $y_i=0$ 的代理人。足够地接近 0 而值得授权的理想点集合满足下面的不等式:

$$h(-\|y_i\|) \leq u'_0$$

由于 $h(\cdot)$ 是严格递增的和连续的,所以对 u'_0 的任意值,集合 $\{y: h(-\|y\|) \leq u'_0\}$ 是闭球。

由于 $h(\cdot)$ 上没有施加任何限制条件,所以,上述结论指出,在空间背景中,授权的信息基础只决定于防范差的结果的意愿;风险偏好、维数和不确定性的性质等都是次要问题。容易看出,放松完全能力假定不会使模型结论有质的变化。随着代理人能力下降,授权集缩小——如果委托人要授权给某个人,而此人对任何事情的了解不比委托人多,则这个代理人的偏好必须十分接近委托人。

如果有多个代理人的理想点落在授权集中,则授权给谁的问题可能更加复杂。传统文献常强调同盟原则:委托人如果授权,就会授权给其偏好最接近自己的代理人。在能力完全相同(即对所有的 i , q_i 相同)时,如果 $x = p - \varepsilon$ (Epstein-O'Halloran 假定的多维度形式),则同盟原则成立。我们把这一结论的证明留做习题。

如果 q_i 彼此不同或者政策结果函数更为一般,联盟原则未必行得通。我们来看代理人在能力上有差异的一个例子。设 $x = p - \varepsilon$, 有两个代理人; $\|y_1\| < \|y_2\|$ 。委托人会授权给代理人 2 而不是更接近自己的代理人 1 吗? Bendor 和 Meirowitz 指出,如果 $q_2 > q_1$, 这种可能性就存在。他们的结论是,如果一个代理人劣于另一代理人——代理人 j 劣于代理人 i , 如果 $\|y_i\| \leq \|y_j\|$ 且 $q_i \geq q_j$, 并且这两个不等式中有一个严格成立——委托人永远不会选择劣代理人。一个更复杂的结论是,一旦放松假定 $x = p - \varepsilon$, 即使 $q_1 = q_2$, 代理人 2 仍有可能被选中。这种结果可能出现,因为不同政策上的不确定性未必相同。在一般模型中,缺少信息的代理人 1 最偏好的政策 p_1 , 可能比缺少信息的代理人 2 最偏好的政策 p_2 , 导致更大的不确定性。如果某些结果(如远离现状的结果)比其他结果(如接近现状的结果)更难实现或者有更大的执行偏差,这种情况就会发生。假设政策空间和结果空间都是 $\mathbb{R}$, 两个代理人的理想点分别为 $y_1 = -1 + \delta$ 和 $y_2 = 1$ 。就是说,代理人 1 更接近委托人。假设 $F(x|p)$ 有如下特征:在 $p > 0$ 时,以相同的概率 $x = p + 0.1$ 或 $x = p - 0.1$; 在 $p < 0$ 时,以相同的概率 $x = p + 0.2$ 或 $x = p - 0.2$ 。在这种情况下,两个代理人如果不知道所发生的冲击,都会选择 $p = y_i$ 。因此,代理人 1 的缺乏信息的政策选择包含更大的结果风险。此时,如果 $h(\cdot)$ 为严格凹函数并且 δ 足够小,委托人就愿意授权给代理人 2, 即使和代理人 1 相比,他的理想点距离委托人更远。习题中给出了这方面的一个例子。

当信息的获得需要付出成本时,代理人之间的搭便车行为可能导致联盟原则不起作用。假设任何代理人或者委托人能以成本 c 了解到冲击。就是说,任何人,只要愿意发生这一成本,就能以 1 的概率观察到所发生的冲击。进一步假设委托

人选择了一位代理人,然后观察他是否支付成本 c 以了解所发生的冲击。如果他不去了解,委托人就收回控制权,并决定自己是否支付成本 c 以了解冲击,并选择政策。委托人也可以不去了解冲击而直接选择政策。在这一背景下,授权集是个多维度的“炸面包圈”。如果委托人授权给理想点与自己接近的代理人,此人可以选择发生成本 c 以得到结果效用 $h(0)$,也可以选择投资,尽管他知道,这样做会导致委托人收回控制权并自己付出成本 c 了解冲击。就是说,了解冲击导致效用为 $h(0) - c$; 而搭便车产生的效用为 $h(-\|y_i\|)$ 。所以,对理想点更接近 $h^{-1}(h(0) - c)$ 的代理人,搭乘委托人的便车是更好的选择。于是,委托人不会授权给在政策空间中十分接近自己的代理人。当然了,位置遥远的代理人所选择的结果不会被委托人接受,就是说,不在授权集中。

命题 11.5 如果获取信息要发生成本 c 并且如果所选中的代理人不发生这一成本,委托人能够收回控制权的话,授权集就由初始授权集中理想点到 0 的距离大于 $d = h^{-1}(h(0) - c)$ 的代理人构成。

Bendor 和 Meirowitz 还在多个代理人和一个委托人背景中探讨了代理人之间的竞争所发挥的作用。假设代理人有完全能力并同时宣布在 X 中的结果。然后,委托人选择一位代理人,由他(他知道所发生的冲击)选择政策以实现他所宣布的结果。如果结果空间是单维度的并且在委托人的两侧都分布着代理人,这一博弈就类似道恩斯竞争。在每个均衡中,至少有两个代理人承诺实现委托人的理想结果。无论政策空间的维数为多少,这一结论都成立。

定义 11.6 (1) 如果不存在向量 $s \in X$ 使得对所有的 $i \in N$, 存在某个 $\lambda_i \in \mathbb{R}$ 使得 $x_i = \lambda_i s$; 或者 (2) 如果存在这样的向量,则肯定存在两个代理人 i 和 j 使得 $\lambda_i > 0$ 且 $\lambda_j < 0$, 我们就说偏好满足多样性。

偏好的多样性要求偏好不是共线性的,如果这个条件得不到满足,则要求在委托人任何一侧,都分布有代理人。

命题 11.6 如果偏好满足多样性,则在每个均衡中至少有两个代理人承诺实现结果 $x = 0$, 委托人从他们中间选择一位。

如果有一位代理人承诺实现 $x = 0$, 则其他代理人的行动显然与收益没有任何关系。所以,在任何策略组合中,如果至少有两位代理人做出这种承诺,这一策略组合就构成均衡。假设没有人做出这种承诺。在博弈中至少有两位代理人时,至少其中的一位能够通过承诺更接近委托人理想点的结果,使最终结果更接近自己的理想点。因此,在任何均衡中,委托人都能得到自己的理想结果。

11.7.5 应用:选举*

再回到选举问题上。这里所探讨的模型比前面分析过的两元选择更接近 Ferejohn 的道德风险模型。在每个时期 $t = 1, 2, \dots$ 中,在职官员首先观察到冲击 $\theta_t \in [z, 1+z]$, θ_t 是他的私人信息,其中 $z \in (0, 1)$ 为常数。然后,他选择某一努

力程度 $a_t \in [0, 1]$ 。每个时期的冲击 θ_t 独立抽取自 $[z, 1+z]$ 上的均匀分布。努力是需要付出成本的: 努力程度 a_t 给官员造成负效用 ca_t 。官员在时期 t 从公职中得到收益 b 。代表性选民关心产出 $x_t = \theta_t a_t$; 产出越高越好。在时期 t , 选民得到收益 x_t , 但是他观察不到 θ_t 或 a_t 。他要决定是让官员留任还是撤换他。在位官员如果被撤换, 则在未来所有时期中, 收益为 0。由于这个模型中只有道德风险而没有逆选择, 所以, 所有可能的官员都有相同的收益, θ 的分布对所有官员都相同。官员和选民使用相同贴现因子 δ 。Ferejohn 指出, 如果选民使用纯策略, 则最优均衡具有截点特征——如果 x_t 大于某个临界值 x^* , 就让在位官员留任; 如果 x_t 小于 x^* , 就撤换他。在这种形式的均衡中, 如果 θ_t 小于某个临界值 θ^* , 在位官员选择 $a_t = 0$; 如果 θ_t 大于 θ^* , 就选择 $a_t = \frac{x^*}{\theta_t}$ 。

但是, 利用激励一致性概念, 我们能构造出混合策略均衡, 和 Ferejohn 的均衡相比, 它能提高选民的福利。假设选民在留任和撤换之间随机做出选择, 留任概率是所观察到的产出 x_t 的函数。那么, 选民能够找到一个留任概率, 它能激励官员选择最高的努力程度吗? 如果在均衡中所有官员都选择相同的努力水平, 则选民在留任和撤换两个决策之间无差异。因此, 随机化构成最优反应。设 $p(x)$ 为在位官员被留任的概率, 它是 x 的函数。给定这一函数, 在时期 t 当在位官员观察到 θ_t 时, 选择 a_t 带给他的期望收益为:

$$p(\theta, a_t) \delta V^{in} - ca_t$$

其中, V^{in} 是在下个时期初这位官员依然在位的博弈所实现的期望效用。在均衡中, 如果此人总是竭尽全力地工作(不管 θ_t 为多少, $a_t = 1$), 则 V^{in} 满足递归关系, 即:

$$V^{in} = b - c + V^{in} \int_a^{1+a} p(\theta) d\theta$$

求解这一方程, 得到:

$$V^{in} = \frac{b - c}{1 - \int_a^{1+a} p(\theta) d\theta}$$

这里的 $\int_a^{1+a} p(\theta) d\theta$ 是在 θ 为未知数和 $a_t = 1$ 时, 官员被留任的期望概率。要诱使官员投入最大的努力程度, 即 $a_t = 1$, 激励一致性条件必须得到满足。所以, 必须通过对 $p(\cdot)$ 的选择使得 $a_t = 1$ 最大化 $p(\theta, a_t) \delta V^{in} - ca_t$ 。因此, $a_t = 1$ 必须是以下一阶条件的解:

$$p'(\theta, a_t) \big|_{a_t=1} = \frac{c}{\theta \delta V^{in}} \quad (11.9)$$

代入有关项, 得到:

$$p'(x) = \frac{c \left(1 - \int_a^{1+a} p(\theta) d\theta \right)}{x \delta (b - c)} \quad (11.10)$$

(11.10)式隐含地定义了函数 $p(\cdot)$ 。由于 $\frac{d \ln x}{dx} = \frac{1}{x}$, 所以, $p(x) = \gamma \ln x + k$ 满足方程(11.10)(其中, γ 和 k 为正常数)。在均衡中, 如果不管 θ_i 为多少, $a_i = 1$, 则相对应的 x 的值包含在 $[z, 1+z]$ 中。要使留任概率是真正的概率, 边界条件 $p(z) \geq 0$ 和 $p(1+z) \leq 1$ 还必须得到满足。因此, 诱生 $a_i = 1$ 的混合策略均衡是否存在的问题, 等价于是否存在 γ 和 k 的值使得以下条件得到满足:

$$\begin{aligned}\gamma \ln z + k &\geq 0 \\ \gamma \ln(1+z) + k &\leq 1\end{aligned}\quad (11.11)$$

$$\frac{\gamma}{x} = \frac{c}{x\delta(b-c)}(1-k-\gamma \int_z^{1+z} \ln x dx)$$

(11.11)式的解可能满足 $\gamma \ln z + k = 0$ 。将这个式子代入到(11.11)式中, 只要 $\gamma \ln(1+z) - \gamma \ln z \leq 1$, 第一个和第二个不等式就得到满足。对应着相应的 k 值, γ 的值是正的。求解 γ , 得到:

$$\gamma = \frac{c}{\delta(b-c)}(1 + \gamma \ln z - \gamma \int_z^{1+z} \ln x dx)$$

或者,

$$\gamma = \frac{c}{\delta(b-c) - \ln z + \int_z^{1+z} \ln x dx} \quad (11.12)$$

由于 $az < 1$ 和 $1+z < e$, 所以 $-\ln z$ 为正。由于 $\ln(\cdot)$ 是增函数且 $z < 1$, 所以, $\ln z < \int_z^{1+z} \ln x dx$ 。这就意味着 $-\ln z + \int_z^{1+z} \ln x dx \geq 0$ 。所以, 在 $b > c$ 时, 方程(11.12)的解是正的。事实上, 我们不必如此费力地推导这一混合策略。前面说过, 条件(11.9)是使得 $a_i = 1$ 为最优选择的条件。实际上, 如果努力的上界为 1, γ 未必要正好等于方程(11.12)的解。我们的解指出, 大于 1 的 a_i 劣于 $a_i = 1$ 。作为练习, 你可以放松约束条件 $a_i \in [0, 1]$ 而假定 $a_i \geq 0$, 找出支持如下均衡的条件: 存在某个固定值 $\alpha > 0$, 使得 $a_i = \alpha$ 。

11.8 机制设计和信号揭示博弈

和很多文献中的处理方式一样, 我们分别介绍了机制设计理论和信号揭示理论。但是, 我们认为, 通过探讨这两个相互独立的专题之间的联系, 能够得到一些启示。

假设信号发送者(参与者 s)的类型为 $\theta \in \Theta = [0, 1]$, 信号空间为 $M = [0, 1]$ 。 θ 服从 Θ 上的分布 $F(\cdot)$, 这是公共知识。信号接收者(参与者 r) 在收到 s 发出的信号 $m \in M$ 后, 选择政策 $p \in X = [0, 1]$ 。甚至在给出收益函数时, 我们依然能够利用激励一致性条件确定发送者的信号充分揭示其类型(即

$m^{-1}(m(\theta)) = \theta$) 的均衡存在的必要条件。在 $m(\theta)$ 为单射的均衡中, 一致的推断肯定集中在正确的 θ 上。因此, 我们可以用如下概率分布表示推断:

$$B(\theta | m') = \begin{cases} 1, & \text{如果 } \theta \geq m^{-1}(m') \\ 0, & \text{在其他情况中} \end{cases}$$

给定充分揭示信息的信号, r 的序贯理性要求:

$$p(m') \in P(m') \equiv \operatorname{argmax}_p u_r(p, m^{-1}(m'))$$

s 的序贯理性要求他没有动机在自己的类型为 θ'' 时, 把自己伪装成类型 θ' , 从而误导 r 。给定映射 $p(\cdot)$, 这一激励一致性条件是, 对所有的 $\theta', \theta'' \in \Theta$,

$$u_s(p(\theta''), \theta'') \geq u_s(p(\theta'), \theta'')$$

接受方的最优反应则要求对所有的 $\theta', \theta'' \in \Theta$,

$$u_r(p_r(\theta''), \theta'') \geq u_r(p_r(\theta'), \theta'')$$

因此, 分离均衡要求 $p(m')$ 同时最大化 $u_s(p, m^{-1}(m'))$ 和 $u_r(p, m^{-1}(m'))$ 。

命题 11.7 只有在参与者的偏好类似——具体地说, 对每个 $\theta \in \Theta$, $\operatorname{argmax}_{p \in X} u_r(p, \theta) \cap \operatorname{argmax}_{p \in X} u_s(p, \theta)$ 非空时, 才存在分离的精炼贝叶斯均衡。

根据这一结论, 在只有一位信号发送者的“空谈”信号揭示模型中, 诚实的信息揭示要求发送者和接收者的收益函数十分相似。我们以第 8 章中开放规则下的 Gilligan-Krehbiel(1989)模型为例, 说明这一结论。在这个模型中, 我们说过, F 对揭示真实类型的信号的最优反应是 $p(-\theta) = \theta$ 和 $p(-\theta) = \theta$ 。命题 11.7 指出, 分离均衡存在当且仅当:

$$u_F(\theta | -\theta) \geq u_F(-\theta | -\theta)$$

$$u_F(-\theta | \theta) \geq u_F(\theta | \theta)$$

$$u_C(\theta | -\theta) \geq u_C(-\theta | -\theta)$$

$$u_C(-\theta | \theta) \geq u_C(\theta | \theta)$$

我们知道, 前两个不等式肯定成立, 因为 $p(-\theta) = \theta$ 且 $p(-\theta) = \theta$, 所以关键的条件是 $-c^2 \geq -(2\theta + c)^2$ 和 $-c^2 \geq -(2\theta - c)^2$ 。如果 $c < \theta$, 这两个不等式都成立。而这正是前面得到的偏好分歧条件。

现在扩展前述一般模型。假设有两个信号发送者, 1 和 2, 他们都观察到 θ , 但是可能有不同的偏好。对是否存在精炼贝叶斯均衡且在这一均衡中信号接收者了解 θ 的问题, 可以看作是个机制设计问题。在这种模型中, 信号接收方对映射 $p(m_1, m_2): M^2 \rightarrow X$ 的选择, 类似对机制的选择。但是与机制设计不同, 信号揭示博弈中的精炼贝叶斯均衡要求接收方的决策在给定的一致推断下, 必须满足序贯理性。这就是说, 我们不能随意选择只满足发送者的激励一致性条件的某种机制。Baron 和 Meiorowitz(2004)指出, 在许多情况中, 信号接收者的行动必须满足序贯

理性的约束条件不是紧的。

在机制设计问题中,假设信号接收者想通过惩罚信号发送者——在他们发送的信号不一致时,惩罚他们——诱使他们说真话。再假设存在着一个差的政策 p^b , 和信号接收者对任何一对说真话的信号做出的最优反应相比,这一政策对两个信号发送者更不利。正式地说,对每个 θ , p^b 使得:

$$\begin{aligned} u_1(\arg\max_{p \in X} u_r(p, \theta), \theta) &\geq u_1(p^b, \theta) \\ u_2(\arg\max_{p \in X} u_r(p, \theta), \theta) &\geq u_2(p^b, \theta) \end{aligned} \quad (11.13)$$

根据 p^b 的定义,下面的机制是两个信号发送者说真话的激励一致性条件:

$$p(m_1, m_2) = \begin{cases} \arg\max_{p \in X} u_r(p, m), & \text{如果 } m_1 = m_2 \\ p^b, & \text{在其他情况中} \end{cases}$$

给定这一政策函数,条件(11.13)指出,对信号发送者1的说真话行为,发送者2的最优反应也是说真话,因为只有这样,才能够避免更糟糕的政策 p^b 。同样的道理,发送者1也有动机说真话。在机制设计框架中,无论信号接收者从 p^b 中得到的效用为多少,仅仅依靠 p^b 的存在,就足以诱使信号发送者说真话。但是,在信号揭示模型中,我们必须探讨选择 p^b 对信号接收者而言是否满足序贯理性。对信号接收者来说,给定他在 $m_1 \neq m_2$ 的所有信息集上形成的推断, p^b 肯定是最优政策。这意味着,对那些熟悉信号揭示模型的人来说,我们的机制设计模型不是很有说服力。但是,弱一致性条件只对以正的概率发生的信息集上的推断施加约束。在信号发送者说真话的均衡中, $m_1 \neq m_2$ 。就是说,满足序贯理性和在 $m_1 \neq m_2$ 时执行政策 p^b 的威胁,并不是对信号揭示模型的根本偏离。它所要求的仅仅是存在 Θ 上的某个分布 $b^b(\cdot)$,使得 $p^b \in \arg\max_p \int u_r(p, \theta) db^b(\theta)$ 。增加一个状态 θ^b ——在这个状态中, $p^b \in \arg\max_p u_r(p, \theta^b)$ ——就足矣。利用这一思路,Baron 和 Meirowitz (2004)得到了下面的结论:

定理 11.4 假设两个信号发送者观察到 θ 。只要存在满足条件(11.13)的政策 $p^b \in X$ 和 Θ 上的分布 $b^b(\cdot)$,使得 $p^b \in \arg\max_p \int u_r(p, \theta) db^b(\theta)$,就存在说真话的精炼贝叶斯均衡。

Gilligan 和 Krehbiel(1989)探讨的是异质立法委员会。Krishna 和 Morgan (2001)则指出,在有两个信号发送者时,存在精炼贝叶斯均衡,信号接收者了解到状态 θ 。这一均衡并不依赖惩罚政策 p^b ,而是利用了非均衡路径上的推断,为惩罚说谎者的政策,提供理性基础。Krehbiel(2001)批判了这一模型,认为对某些信号的非均衡路径反应,是高度不连续的,并且向错误的方向变化。虽然在均衡中,高政策是对低信号的最优反应,但是,Krishna 和 Morgan 的精炼贝叶斯均衡要求,面对着非均衡路径上的高信号时,使用低政策。

Battaglini(2002)指出,如果政策空间和状态空间都是多维的(例如, $X = \Theta =$

$[0, 1]^2$), 则能构造出说真话均衡, 它不以上述怪异方式决定于推断。假设有一位信号接收者和两位信号发送者, 他们有两维空间上的欧几里得偏好。信号接收者的理想点为 $(0, 0)$, 发送者 1 的理想点为 $(1, 0)$, 发送者 2 的理想点为 $(0, 1)$ 。在这个模型中, 信号发送者都能完全观察到冲击 $\theta \in \Theta$, 结果 x 是政策 p 的随机函数, 设 $x = p + \theta$ 。对 $s \in \{1, 2\}$, 信号为 $m_s = (m_s^1, m_s^2)$ 。结果和政策都为向量: $x = (x^1, x^2)$, $p = (p^1, p^2)$ 。我们可以构造出一个相当简单的直接机制。虽然两个发送者的偏好都不同于信号接收者^①, 但是, 每个发送者分别在其中一个维度上有和接收者相同的偏好。发送者 1 和接收者都希望结果向量中第二个坐标值尽可能接近 0, 发送者 2 和接收者都希望结果向量中第一个坐标值尽可能接近 0。假设理想点为 $(1, 0)$ 的发送者 1 能完全控制第二个坐标, 就有 $p^1 = -m_1^1$; 假设理想点为 $(0, 1)$ 的发送者 2 能完全控制第一个坐标, 就有 $p^2 = -m_2^2$ 。给定这一机制, 信号 $m_s(\theta) = \theta$ 是最优反应, 这一机制因此满足激励一致性。那么, 接收方所使用的策略能否得到精炼贝叶斯均衡的支持呢? 如果它能够得到均衡支持的话, 我们就需要找到弱一致的推断, 使得给定这一推断, 政策函数满足序贯理性。由于前述映射 $p(m)$ 在所到达的每个信息集上都是直接机制, 所以, 根据贝叶斯规则, 推断为 $\theta = m_1 = m_2$ 上。我们还得为 $m_1 \neq m_2$ 的信息集找到推断, 使得给定这一推断, 对信号接收者来说, 政策函数 $p(m)$ 满足序贯理性, 对信号发送者来说, 说真话满足序贯理性。一种简单的处理方法是考虑 m_1^2 和 m_2^1 而形成推断。换句话说, 推断集中在 $\theta = (m_1^2, m_2^1)$ 上。Battaglini 指出, 如果偏好为多维欧几里得偏好, 如果有两个信号发送者, 只要三个参与者的理想点不在一条直线上, 就存在分离精炼贝叶斯均衡。

我们看到, 在某些偏好类型下, 在有两个信号发送者时, 存在分离精炼贝叶斯均衡。在有三个拥有完全信息的信号发送者 $\{1, 2, 3\}$ 时, 我们不需要对偏好施加任何假设就能找到一个简单的分离精炼贝叶斯均衡。设 $p^*(\theta)$ 是取自对应 $\operatorname{argmax}_x u_r(p, \theta)$ 的一个选择(selection), 假设对某个 θ , $p^+ \in p^*(\theta)$ 。和前面的例子一样, 假设信号接收者能使用下面的机制, 它依赖于信号 m_1 、 m_2 和 m_3 :

$$p(m) = \begin{cases} p^*(\theta), & \text{存在 } i, j \in \{1, 2, 3\}, \text{ 使得 } \theta = m_i = m_j \\ p^+, & \text{在其他情况中} \end{cases}$$

在这里, 说真话构成最优反应, 因为如果 i 和 j 在说真话, 则 k 的信号对结果没有任何影响。容易证明, 这一政策函数满足序贯理性。在 $m_1 = m_2 = m_3 = \theta$ 时, 集中在 θ 上的任何推断, 都满足弱一致性; 对 $i, j \in \{1, 2, 3\}$, 在 $\theta = m_i = m_j$ 时, 集中在 θ 上的任何推断, 都使得上述政策映射构成信号接收者的最优反应。

Baron 和 Meirowitz 的结论是, 任何直接机制, (1) 在接收到说真话的信号后, 为接收者选择了最优政策; 并且 (2) 在接收到虚假信号后, 选择了某个政策, 且给定

① 根据命题 11.7, 在只有一位信号发送者的博弈中, 不存在说真话均衡。

⑥上的某个推断,它是最优政策,则在信号揭示博弈中,这个机制能得到某个精炼贝叶斯均衡的支持。

11.9 习题

习题 11.1 假设系主任想最大化自己的剩余(总出资额减去咖啡机支出),而不是最大化整个政治学系的福利。他依然想实行 Groves-Clarke 机制吗?

习题 11.2 在一个机制设计问题中,有两个代理人。机制设计者选择政策 $p \in [0, 1]$, 以求最大化代理人 1 和代理人 2 的收益之和。每个代理人的收益为 $u_i(p, y_i, \beta_i) = -\beta_i |y_i - p|$ 。代理人类型 (y_i, β_i) 中的每个坐标服从均匀分布。假设它们相互独立。(1)设计一个直接机制;(2)存在直接机制能(在贝叶斯 Nash 均衡中)执行设计者的目标吗?

习题 11.3 这道习题要求证明另一形式的显示原理。选择函数 $g(\theta)$ 是在占优策略中可执行的,如果存在机制 $\langle M, p(\cdot) \rangle$, 其中,每个人都有一个占优策略,在给定政策函数 $p(m(\theta)) = g(\theta)$ 时,每个 $\theta \in \Theta$ 使用策略 $m_i(\theta)$ 。证明:给定机制设计问题 $\langle \Theta, F(\theta), u \rangle$, 存在直接机制 $\langle \Theta, p(\cdot) \rangle$, 如果选择函数 $g(\theta)$ 是在占优策略上可执行的,则这一机制在占优策略上执行 $g(\theta)$ 。

习题 11.4 证明 Moulin 的结论(命题 11.2)。

习题 11.5 假设有一群代理人,每个人都有二次方偏好,并且每个人拥有关于自己理想点的私人信息。假设他们被要求报告自己的理想点,然后执行政策 $x(m) = \frac{1}{n} \sum m_i$ (即所执行的政策为平均理想点)。证明在这一博弈中,说真话并不构成贝叶斯 Nash 均衡。

习题 11.6 证明(11.1)式所描述的策略构成最高价格拍卖的均衡。

习题 11.7 证明:“举起号码牌直至价格超过 $b(\theta)$ ”的策略,构成价格降序拍卖的精炼贝叶斯均衡。

习题 11.8 如果 $F(\cdot)$ 是 $[0, 1]$ 上的均匀分布,求标准拍卖的期望收入。

习题 11.9 我们来探讨存在道德风险的 Ferejohn 模型。假设(11.8)式并不成立。构造混合策略均衡,在这一均衡中,政府有时候选择 $a = \text{低}$; 在 $x = \text{低}$ 时,选民总是撤换政府,在 $x = \text{高}$ 时,选民偶尔也会撤换政府。

习题 11.10 证明:在 Epstein-O'Halloran 模型中, P^* 是个闭区间 $[p, \bar{p}] \subset X$ 。

习题 11.11 在 Epstein-O'Halloran 模型中,证明在多数通过规则下得到的结果 P^* , 是理想点为 l_m 的立法者偏好的结果。

习题 11.12 在 Epstein-O'Halloran 模型中,假设 ε 服从分布 $N(0, \sigma^2)$ 。计算最优政策 P , P 如何决定于 l_m 、 a 和 σ^2 ? 提示: $E[\varepsilon | \varepsilon < m] = -\frac{\sigma\varphi\left(\frac{m}{\sigma}\right)}{\Phi\left(\frac{m}{\sigma}\right)}$ 。

习题 11.13 这道习题取自 McCarty(2004)。在 Epstein-O'Halloran 模型中, 假设理想点为 $g > l_m$ 的总统 G , 在 L 选择 P 之前, 任命 A (就是说, 选择了 a)。假设 A 了解 ε , 而 G 和 L 认为 ε 在 $[-E, E]$ 上均匀分布。证明总统的最优任命应是 $a^* \in (l_m, g)$ 。 a^* 如何决定于 l_m 和 E ? 如果总统在 L 选择 P 之后任命 A 呢?

习题 11.14 在 Bendor-Meirowitz 模型中, 证明 $\{y: h(\|y\|) \leq u'_0\}$ 是闭球。

习题 11.15 在 Bendor-Meirowitz 模型中, 如果对所有的 i, q_i 相同, 且 $x = p - \varepsilon$ (ε 服从 $\mathbb{R}^d$ 上的分布 $F(\cdot)$), 则联盟原则成立。

习题 11.16 在 Bendor-Meirowitz 模型中, 假设 $q_1 = q_2 = \frac{3}{4}$, $u(x) = -x^2$ 。找出使得委托人愿意授权给代理人 2 的最小的 δ 值。

习题 11.17 在连续的 Ferejohn 模型中, 假设 a_i 是任意非负的努力程度。请把 x^* 和 θ^* 的均衡值表示为外生参数 b, c, z 和 δ 的函数。

习题 11.18 在连续的 Ferejohn 模型中, 假设 a_i 为正数 (可以大于 1)。对给定的 $\alpha > 0$, 找出一个均衡, 在每个时期里无论 θ_i 为多少, $a_i = \alpha$ 。

习题 11.19 在 Battaglini 模型中, 设信号发送者 1 的理想点为 $(1, 1)$, 发送者 2 的理想点为 $(0, 1)$ 。请构造完全分离的精炼贝叶斯均衡。

习题 11.20 假设政策制定者必须选择政策 $p \in \mathbb{R}^3$ 。政策结果为 $x = p + \varepsilon$, 其中 ε 为随机冲击。政策制定者的先验推断是, ε 服从分布函数 $F(\cdot)$ 。假设有五个代理人。代理人 1、代理人 2 和代理人 3 都能完全观察到冲击 ε 的第一个坐标 (ε^1), 但是观察不到其余坐标。代理人 4 和代理人 5 能够观察到冲击 ε 的第二个和第三个坐标 ($\varepsilon^2, \varepsilon^3$)。假设政策制定者的理想点是向量 0 。每个代理人的理想点为 $\mathbb{R}^3$ 中任意点。所有的代理人有结果 x 上的二次方偏好。在这一问题中, 是否存在 “cheap talk” 博弈, 它有精炼贝叶斯均衡, 均衡结果始终为政策制定者最偏好的结果 0 ? 如果存在这样的博弈, 请给出个例子, 并找出一个均衡。如果不存在这样的博弈, 就请给出证明。



数学附录

数学是对公理体系和这些体系的真命题的研究。数学是种语言,方便我们分析这些体系和最终得到的命题。和日常语言相比,数学陈述能更明确地揭示出假设和结论之间的关系。数学命题如果得自产生这些命题的公理体系中,则为“真”或者说是正确的。例如,欧几里得几何是由平行线不相交这一公理基础上的所有真命题构成的集合。两点之间直线距离最短的命题在欧几里得几何中是真命题。但是,从纽约中央公园到普林斯顿的最快捷的路线是直线的命题未必成立,因为这决定于这两个点之间的交通条件、路面质量以及其他因素。我们并不清楚有关距离的数学上的真命题在多大程度上准确描述了我们所关心的实际问题,同样地,博弈模型中的在数学上正确的命题,也未必能准确地反映现实政治。想了解真相而不仅仅追求模型的价值和数理形式的学生和学者,应该寻求认识论学家的帮助。在我们看来,关键是确定在数学上正确的命题究竟是什么。在这部分中,我们首先给出数学术语的明确定义,然后给出在政治博弈论中有着大量应用的数学体系的标准公理,接着给出政治博弈论中用到的重要数学命题。

12.1 数学命题和证明

和其他任何语言一样,数学的基本构成单位是命题。我们用 P 表示命题。以下是读者可能遇到的各种数学命题。

- 全称命题:在给定的数学系统内, P 始终为真。

我们来看一个全称命题。设 x 是个实数, $\forall x, x \leq |x|$ 。其中, $\forall$ 表示“对所有的”或者“对每一个”。要证明这种命题,我们要求只使用 x 的每个值都具有的共同性质,证明对任意的 x ,这一命题成立。

- 存在命题:存在一些条件,在这些条件下, P 为真。

我们来看一个存在命题。 $\exists x$ 使得 $x = |x|$ 。其中, $\exists$ 表示“存在”。要证明存在命题,我们只需在给定的系统内找到一个 x ,对这个 x , P 为真。在这里,所需的条件是 $x \geq 0$ 。

在数学中,有一些方法可用于证明某一命题为真命题。这里介绍在本书中用到的证明方法。

12.1.1 演绎

在用演绎方法进行的证明中,某个命题为真,是因为它与某个已知为真或者被假定为真的命题存在着逻辑联系。假设我们知道 P 为真命题;如果能证明“如果 P ,则 Q ”,或者说,如果能证明“ $P \Rightarrow Q$ ”,我们即证明了 Q 是真命题。这种证明可能需要多个步骤,如:

$$P \Rightarrow R$$

$$R \Rightarrow S$$

$$S \Rightarrow Q$$

有时在我们揣摩命题的证明思路时,首先想到的可能是证明 $S \Rightarrow Q$ 。但是,在揭示证明的逻辑结构时,却应该按照上述步骤进行。演绎法既可用于证明存在命题也可用于证明全称命题。

12.1.2 逆否命题证明

有时候,为证明 $P \Rightarrow Q$,可首先用否命题 $\sim Q$ 和 $\sim P$ 表述它,其中, $\sim$ 表示“无”。($\sim Q \Rightarrow \sim P \Rightarrow (P \Rightarrow Q)$),这在逻辑上成立。

例题 12.1 如果 $7m$ 是奇数,则 m 是奇数。

在这里, $P = \{7m \text{ 是奇数}\}$, $Q = \{m \text{ 是奇数}\}$, $\sim P = \{7m \text{ 是偶数}\}$, $\sim Q = \{m \text{ 是偶数}\}$ 。我们通过证明 $\sim Q \Rightarrow \sim P$,来证明 $P \Rightarrow Q$ 。

$$\begin{aligned} \sim Q &\Rightarrow \text{存在某个整数 } k \text{ 使得 } m = 2k \\ &\Rightarrow 7m = 7(2k) \\ &\Rightarrow 7m = 2(7k) \\ &\Rightarrow \text{存在某个整数 } n \text{ 使得 } 7m = 7n \\ &\Rightarrow \sim P \end{aligned}$$

12.1.3 反证法

证明命题 P 为真命题的另一种方法是证明 $\sim P$ 是假的。要证明一个命题是假的,我们只需找到一个反例。只要有一个反例证明 $\sim P$ 是假的,就足以证明 P 是真命题。在否定一个全称命题时或者在证明一个存在命题为真时,这种方法很管用。

例题 12.2 设 n 是任意整数, $P = \{\text{存在 } n > 0 \text{ 使得 } n^2 + n + 17 \text{ 不是素数}\}$ 或者 $P = \{\exists n > 0 \ni n^2 + n + 17 \text{ 不是素数}\}$,其中, $\exists$ 表示“存在”, $\ni$ 表示“使得”。

构造 $\sim P = \{\forall n > 0, n^2 + n + 17 \text{ 是素数}\}$ 。 $\sim P$ 是假的,因为存在 $n = 17$,使得

$n^2 + n + 17 = 17 \cdot 19$, 因此, $n^2 + n + 17$ 不是素数。所以, P 是真的。

我们还可以从 $\neg P$ 出发, 通过一系列推导步骤, 得到一个假命题, 借此证明 $\neg P$ 是假的。

$$\neg P \Rightarrow \text{对所有的 } n, n+1+\frac{17}{n} \text{ 不是整数}$$

$$\Rightarrow \text{对所有的 } n, \frac{17}{n} \text{ 不是整数}$$

$$\Rightarrow 1 \text{ 不是整数}$$

其中的最后一个命题显然是假的。

12.1.4 数学归纳法

一些命题描述的是指数 n (整数) 的性质, 我们可以把这一性质写为 $P(n)$ 。证明命题 $P(n)$ 对每个自然数都成立的一种方法, 是证明 $P(1)$ 为真, 且如果 $P(n)$ 为真, 则 $P(n+1)$ 为真。

例题 12.3 设 $P(n)$ 说的是, $2n < n^2$ 。对 $n \geq 3$, 求证 $P(n)$ 为真。在 $n=3$ 时, $6 < 9$, 因此, $P(3)$ 为真。假定 $P(n)$ 为真, 即 $2n < n^2$ 。由于 $2(n+1) = 2n+2$ 且 $(n+1)^2 = n^2 + 2n + 1$, 这一假定意味着如果 $2 < 2n+1$, 则 $P(n+1)$ 为真。这一条件可简化为 $\frac{1}{2} < n$ 。对 $n > 3$, 它显然为真。

12.2 集合和函数

12.2.1 集合

集合是不同事物构成的一个全体。集合中的事物被称为元素。我们可以通过列举集合中的元素或描述集合中元素的特征, 来表示集合。假设 S 是小于 10 的所有正整数构成的集合, 我们可以把它写为 $S = \{1, 2, 3, 4, 5, 6, 7, 8, 9\}$ 或者写为 $S = \{n \mid n \text{ 是整数且 } 0 < n < 10\}$ 。其中的第二种表述读为“ S 是所有大于 0 和小于 10 的整数构成的集合”。集合或者有有限个元素, 如这里的集合 S ; 或者有无限个元素, 如 $I = \{n \mid n \text{ 是大于 7 的整数}\}$ 或者如 $J = \{x \mid 0 < x < 1\}$ 。无限集可以是可数集也可能是不可数集。在这里, 集合 I 是可数集, 因为我们能把 I 中的每个元素与不同的整数对应起来。 J 是不可数集, 因为我们无法把它的每个元素与整数对应起来。

例题 12.4 有理数是可以表示为分数 $\frac{p}{q}$ 的数, 其中 p 和 q 都是整数。证明:

$\mathbb{Q} = \{x \mid x \text{ 是有理数}\}$ 是个无限的但可数的集合。

我们用符号“ $\in$ ”表示某个元素属于一个集合。于是, 以下陈述成立: $3 \in S$,

$9 \in I, 0.345678 \in J, 0.8 \in Q$ 。我们用符号“ $\notin$ ”表示不属于某个集合的元素,例如 $10 \notin S, 7 \notin I, 0 \notin J, \pi \notin Q$ 。

12.2.2 集合关系和集合运算

下面给出了不同集合之间的一些有用的关系。

● 集合相等

$S_1 = S_2$ 意味着 $x \in S_1$ 当且仅当 $x \in S_2$ 。换句话说,如果 $S_1 = S_2$, 则 S_1 和 S_2 包含完全相同的元素。

● 子集

如果 $x \in S_1$ 意味着 $x \in S_2$, 则说 S_1 是 S_2 的子集或者 $S_1 \subseteq S_2$ 。就是说,如果 S_1 中的所有元素也是 S_2 中的元素,则 S_1 是 S_2 的子集。如果 $S_1 = S_2$, 则 $S_1 \subseteq S_2$ 且 $S_2 \subseteq S_1$ 。我们还可以定义“真”子集,它排除了两个集合相等的可能性。 $S_1 \subset S_2$ 意味着对所有的 $x \in S_1, x \in S_2$, 且存在 $y \in S_2$ 使得 $y \notin S_1$ 。换句话说, S_1 中的所有元素都在 S_2 中,但是 S_2 还包含有 S_1 中没有的其他元素。我们用符号 $\supset$ 和 $\subset$ 表示以相反顺序概括的集合之间的上述关系。符号“ $\not\subseteq$ ”表示“…不是…的子集”。

● 不交

两个集合如果没有共同元素,则说它们不交。正式地说,如果 $x \in S_1$ 意味着 $x \notin S_2$ 且 $x \in S_2$ 意味着 $x \notin S_1$, 则 S_1 和 S_2 不交。

零集或空集是个特殊的集合,常写为 $\emptyset$,是如下定义的: $\emptyset$ 不包含任何元素,并且对所有的集合 $S, \emptyset \subset S$ 。 $\emptyset$ 显然是最小的集合。我们有时也会用到集合 U , 称其为全集。对我们所讨论的所有的集合 $S, S \subseteq U$ 。

在集合上有几种数学运算。

● 并

两个或多个集合的并,是由至少属于其中一个集合的所有元素构成的集合。正式地说, S_1 和 S_2 的并为 $S_1 \cup S_2 = \{x \mid x \in S_1 \text{ 或 } x \in S_2\}$ 。用 i 标示的一组集合的并为 $\bigcup_{i=1}^n S_i = S_1 \cup S_2 \cup \cdots \cup S_n$ 。对于用指数集 I 标示的无限集合族的并,可以表示为 $\bigcup_{i \in I} S_i$ 。

● 交

两个或多个集合的交,是所有这些集合所共有的元素构成的集合。正式地说, S_1 和 S_2 的交为 $S_1 \cap S_2 = \{x \mid x \in S_1, x \in S_2\}$ 。用 i 标示的一组集合的交为 $\bigcap_{i=1}^n S_i = S_1 \cap S_2 \cap \cdots \cap S_n$ 。如果 S_1 和 S_2 不交,则 $S_1 \cap S_2 = \emptyset$ 。由于这两个集合没有任何共同的元素,所以这两个集合的交的唯一子集是零集。对于用指数集 I 标示的无限集合族的交,可以表示为 $\bigcap_{i \in I} S_i$ 。

● 补

给定全集 U 和子集 S, S 的补为 $S^c = U \setminus S = \{x \mid x \in U, x \notin S\}$ 。

● De Morgan 法则

设 A 、 B 、 C 和 D 都是集合, 这些集合上的运算满足以下性质:

交换律: $A \cup B = B \cup A$, $A \cap B = B \cap A$

结合律: $(A \cup B) \cup C = A \cup (B \cup C)$, $(A \cap B) \cap C = A \cap (B \cap C)$

分配律: $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$, $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$ 。

● 积

集合 A 和 B 的积, 被写为 $A \times B$, 是有序数对 (a, b) 构成的集合, 其中, $a \in A$, $b \in B$ 。对于用 i 标示的集合, 我们把它们的积写为 $\prod_{i=1}^n S_i = S_1 \times S_2 \times \cdots \times S_n$ 。有人使用 $\times_{i=1}^n S_i$ 表示 $\prod_{i=1}^n S_i$ 。对于用集合 I 标示的无限集合族, 可以把它们的积写为 $\prod_{i \in I} S_i$ 。

12.2.3 对应和函数

建立两个集合之间关系的一种方法是指出在这两个集合中哪些元素“相对应”或者“结合在一起”。一般而言, 对应是个规则 f , 它把 S_1 的元素和 S_2 的元素联系在一起。在数学上, 我们把对应写为 $f: S_1 \rightarrow S_2$, 其中 S_1 被称为定义域(或前像集), S_2 被称为值域(或像集)。例如, 在图 12.1 中, 对应 f 把 $x \in S_1$ 与 $y, z \in S_2$ 联系在一起。

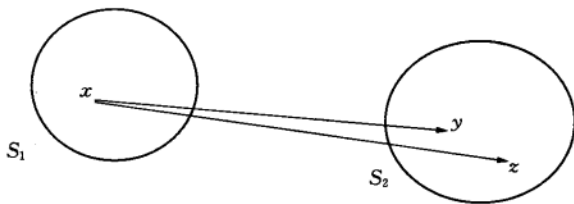


图 12.1 对应

函数是特殊对应, 它把定义域中的每个元素, 与值域中唯一一个点联系在一起。就是说, 如果对应 $f: S_1 \rightarrow S_2$ 是个函数, 则对每个 $x \in S_1$, $f(x)$ 是 S_2 中的一个元素。对函数, 我们将其写为 $f: S_1 \rightarrow S_2$ 。在函数中, 定义域中的多个点可以映射到值域中的同一个点上。在图 12.2 中, 函数把 $x, w, v \in S_1$ 映射到元素 $y, z \in S_2$ 。由于定义域中的每个元素被映射到值域中的唯一一个元素上, 所以, 我们可以将函数写为 $y = f(x)$, 其中, $x \in S_1$, $y \in S_2$ 。我们还可以把函数表示为有序数对集的形式, 例如 $\{(y, x) \mid \text{存在 } x \in S_1, \text{ 使得 } y = f(x)\}$ 。

给定函数 $f: A \rightarrow B$, 我们可以定义它的两个重要性质:

● 单射性: 对所有的 a_1 和 $a_2 \in A$, $f(a_1) = f(a_2)$ 当且仅当 $a_1 = a_2$ 。值域中的每个点与定义域中唯一一个点相对应。这一性质也被称为“一对一”映射。

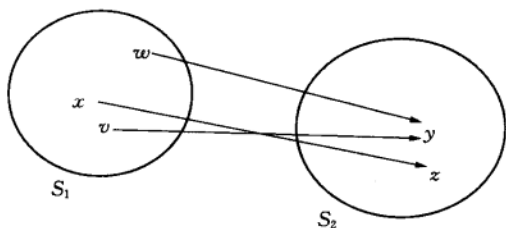


图 12.2 函数

• 满射性: 对每个 $b \in B$, 存在 $a \in A$ 使得 $f(a) = b$. 这一性质也被称为“到上”映射。

一个函数如果同时具有这两个性质, 就被称为双射。此时, 它有反函数, 反函数把 B 中的点映射到 A 中的点。我们将这一反函数写为 $f^{-1}: B \rightarrow A$ 或者 $f^{-1}(b) = a$, 其中, $b \in B, a \in A$ 。

例题 12.5 $y = 2x$ 是双射。由于每个 y 被映射到一个 x 上, 所以, 我们可以把它的反函数写为 $f^{-1}(x) = \frac{y}{2}$ 。

例题 12.6 $y = x^2$ 不是单射, 因为 x 和 $-x$ 对应着同一个 y 。如果我们把定义域限制为正数, 它就有反函数, 可写为 $f^{-1}(x) = \sqrt{y}$ 。

12.3 实数

实数 $\mathbb{R}$ 由所有的整数和有理数(整数的商)与无理数(不是整数的商的数)构成。实数是个具有某些结构的数集。它有两种运算, 即 $+$ 和 $\times$, 它们都是从 $\mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ 的映射; 在实数上, 有弱序 $\geq$, 它是 $\mathbb{R} \times \mathbb{R}$ 的子集。这两种运算就是我们所熟悉的加运算和乘运算, 而这一弱序便是“大于或等于”。实数是被四个公理定义的。就我们的目的而言, 只需要强调其中部分公理。实数系是个域; 就是说, $+$ 和 $\times$ 具有我们在小学阶段所学到的性质; 加(或乘)的先后顺序并不重要, 乘运算服从分配律(就是说, $\forall x, y, z \in \mathbb{R}, x(y+z) = xy + xz$, 与 0 和 1 相乘或相加, 得到的便是我们早就学到的结果, 每个非零实数有乘法逆(即 $x \times \frac{1}{x} = 1$)。而且, 实数满足序公理, 序公理保证了 $\geq$ 具有我们所期望它具有的特征。实数还满足完备性公理, 它可能是你不熟悉的, 我们给出它的正式定义。

定义 12.1 完备性公理: 对每个非空子集 $S \subset \mathbb{R}$, 如果存在 S 的上界 b (就是说, $x \in S \Rightarrow x \leq b$), 则存在最小上界 c (就是说, c 是 S 的上界, 并且, 如果 z 是 S 的上界, 则 $c \leq z$)。换句话说, 每个有上界的集合有最小上界。

$\mathbb{R} \setminus \mathbb{Z}$ 不是完备空间, 这里的 $\mathbb{Z}$ 是整数集。 $(0, 1)$ 在这一空间中有上界 $\left(\frac{3}{2}\right)$ 是

其上界)却没有最小上界。根据这一公理,实数中没有任何窟窿。极限的许多性质依赖于实数的这一性质。

实数列的极限

实数列 $\{x_n\}_{n=1}^{\infty}$, 有时被写为 $\{x_n\}$, 是由实数构成的无穷列。更准确地说, 一个实数列是一个函数, 它把可数数集 $(1, 2, 3, \dots)$ 映射到实数上。在这一意义上, x_n 是这一函数在整数 n 上的取值。当我们抽取整数集的某个子集时, 由此形成的数列便是子数列。

定义 12.2 $l \in \mathbb{R}$ 是数列 $\{x_n\}$ 的极限, 如果对每个 $\varepsilon > 0$, 存在 N 使得对所有的 $n > N$, $|x_n - l| < \varepsilon$ 。如果 l 是数列 $\{x_n\}$ 的极限, 我们就将其写为 $l = \lim x_n$ 。

命题 12.1 一个数列至多有一个极限。

这里给出这一命题的证明, 借此演示如何证明数学陈述。

证明: 反证法。假设不是如此, 设 $a = \lim x_n = b$ 且 $a \neq b$ 。由于 $a \neq b$, 存在某个 $\varepsilon > 0$, 使得 (1) $|a - b| > 2\varepsilon$ 。由于 $a = \lim x_n = b$, 所以肯定存在某个 N , 当 $n > N$ 时, (2) $|x_n - a| < \varepsilon$ 且 (3) $|x_n - b| < \varepsilon$ 。不失一般性, 假设 $a < b$ 。如果 $x_n < a < b$, 结论 (1) 和结论 (3) 相矛盾; 如果 $a < b < x_n$, 结论 (1) 和结论 (2) 相矛盾; 如果 $a < x_n < b$, 则结论 (1) 意味着或者结论 (2) 或者结论 (3) 不成立。这三种情况中必有一种成立。 $\square$

定义 12.3 数列 $\{x_n\}$ 是柯西数列, 如果对每个 $\varepsilon > 0$, 存在 N 使得对所有的 $n, m > N$, $|x_n - x_m| < \varepsilon$ 。

命题 12.2 一个数列有极限当且仅当它是柯西数列。

定义 12.4 数列 $\{x_n\}$ 收敛于无穷大 ∞ (负无穷大 $-\infty$), 如果对任意的 $b \in \mathbb{R}$, 存在某个 N 使得对所有的 $n > N$, $x_n > b$ ($x_n < b$)。

定义 12.5 $l \in \mathbb{R}$ 是数列 $\{x_n\}$ 的聚点 (cluster point), 如果对每个 $\varepsilon > 0$ 和对每个 N , 存在某个 $n > N$, 使得 $|x_n - l| < \varepsilon$ 。

数列 $x_n = (-1)^n$ 有两个聚点, 1 和 -1, 但是没有极限。

命题 12.3 $l \in \mathbb{R}$ 是 $\{x_n\}$ 的聚点当且仅当它是子数列 $\{x_{n_k}\}$ 的极限。

定义 12.6 $l \in \mathbb{R}$ 是数列 $\{x_n\}$ 的上确界 ($\limsup$), 如果 (1) 对每个 $\varepsilon > 0$, 存在 N 使得对所有的 $n > N$, $x_n < l + \varepsilon$; (2) 对每个 $\varepsilon > 0$ 和任意的 N , 存在某个 $n > N$, 使得 $x_n > l - \varepsilon$ 。我们将其表示为 $l = \limsup x_n$ 。

定义 12.7 $l \in \mathbb{R}$ 是数列 $\{x_n\}$ 的下确界 ($\liminf$), 如果 $l = -\limsup(-x_n)$ 。

下面的定义可能更加直观。

定义 12.8 上确界是最大的聚点, 下确界是最小的聚点。

命题 12.4 对任意数列 $\{x_n\}$, $\limsup x_n \geq \liminf x_n$; 并且, 如果等号成立, 则存在 $\lim x_n$ 且 $\lim x_n = \limsup x_n = \liminf x_n$ 。

12.4 点和集合

现在介绍带有一种特定结构的空问,我们称之为度量空问。度量空问上定义了距离函数,它满足以下特征。

定义 12.9 度量空问 (X, d) 由点集 X 和满足以下性质的距离函数 $d(x, y): X \times X \rightarrow \mathbb{R}$ 构成:

- (1) $d(x, y) \geq 0$;
- (2) $d(x, y) = 0$ 当且仅当 $x = y$;
- (3) 对任意的 $x, y \in X, d(x, y) = d(y, x)$;
- (4) 对任意的 $x, y, z \in X, d(x, z) \leq d(x, y) + d(y, z)$ 。

给定任意的距离函数,我们可以定义一个点周围的区域,称之为“球”。对任意实数 $\varepsilon > 0$ 和点 $x \in X$, 围绕着 x 的 ε -球为子集 $B(x, \varepsilon) = \{y \in X: d(x, y) < \varepsilon\}$ 。在度量空问中,集合还有两个重要特征。

定义 12.10 集合 $A \subset X$ 是开的,如果对每个点 $x \in A$, 存在某个 $\varepsilon > 0$ 使得 $B(x, \varepsilon) \subset A$ 。集合 $A \subset X$ 是闭的,如果它的补 $X \setminus A$ 是开的。

集合 X 和 $\emptyset$ 既是开集也是闭集。我们给出有关开集和闭集的几个结论,它们告诉我们如何对开集和闭集进行某些集合运算,并保证通过这些运算得到的结果继承了初始集合的开性和闭性。

命题 12.5

- (1) 如果 O_1 和 O_2 是开的,则 $O_1 \cap O_2$ 是开的。
- (2) 给定开集族 $O_1, O_2, \dots$, 集合 $\bigcup_i O_i$ 是开的。
- (3) 如果 C_1 和 C_2 是闭的,则 $C_1 \cup C_2$ 是闭的。
- (4) 给定闭集族 $C_1, C_2, \dots$, 集合 $\bigcap_i C_i$ 是闭的。

开集的无限交未必是开的。在集合族 $\left(-\frac{1}{n}, \frac{1}{n}\right)$ 中,尽管每个集合都是开集,但是其交集为集合 $\{0\}$, 不是开的。

我们再定义集合的另一个特征。

定义 12.11 集合 $A \subset X$ 是有界的,如果存在某个实数 k , 使得对每个 $x, y \in A$, $d(x, y) < k$ 。

如果 X 是有限维欧几里得空问 $\mathbb{R}^n = \{(x_1, x_2, \dots, x_n): x_i \in \mathbb{R}\}$ 的子集,则有以下定义:

定义 12.12 集合 $A \subset \mathbb{R}^n$ 是紧集,如果 A 是闭的和有界的。

在任意度量空问中,我们需要使用紧性的更一般的定义。这个定义(或者说,紧性的拓扑定义)用到了开覆盖的概念。

定义 12.13 给定集合 A , A 的开覆盖是集合族 $\{O_\theta\}_{\theta \in \Theta}$, 满足 $A \subset \bigcup_{\theta \in \Theta} O_\theta$, 其中 Θ 是任意指数集且对每个 $\theta \in \Theta, O_\theta$ 是开的。换句话说,如果 $x \in A$, 则存在某个 $\theta \in \Theta$, 使得 $x \in O_\theta$ 。

集合是紧的,如果它的每个开覆盖都有有限子覆盖。这一定义的正式形式为:

定义 12.14 集合 A 是紧的,如果给定 A 的任意开覆盖 $\{O_\theta\}_{\theta \in \Theta}$, 存在某个有限集 $B \subseteq \Theta$, 使得 $\{O_\theta\}_{\theta \in B}$ 是 A 的覆盖。

在度量空间中,如果定义了加法运算,如在 $\mathbb{R}^n$ 中,就有下面的定义:

定义 12.15 集合 A 是凸的,如果对每个 $x, y \in A$ 和每个实数 $\lambda \in [0, 1]$, 点 $\lambda x + (1-\lambda)y$ 在 A 中。

下面的结论在社会选择理论和讨价还价理论中常被用到。例如,如果每个人的偏好都是严格凸的,下面的结论就能保证由这些人构成的集体所偏好的点组成的集合(即上半集)是严格凸的。

命题 12.6 给定一个集合族 $\{A_i\}_{i \in \alpha}$, 如果对每个 $i \in \alpha$, 集合 A_i 是凸的, 则 $\bigcap_{i \in \alpha} A_i$ 是凸集。

12.5 函数的连续性

函数的一个重要性质是连续性。设 X 和 Y 是两个度量空间的子集,这些空间上的度量分别为 d_X 和 d_Y 。例如, X 和 Y 可以为实数空间,即 $X=Y=\mathbb{R}$ 。但是,下面的定义和结论适用于把一个度量空间映射到另一个度量空间的任何函数。

定义 12.16 函数 $f: X \rightarrow Y$ 在点 $x \in X$ 上是连续的,如果对任意的 $\varepsilon > 0$, 存在某个 $\delta > 0$, 使得对满足 $d_X(x, y) < \delta$ 的每个 $y \in X$, $d_Y(f(x), f(y)) < \varepsilon$ 。函数是连续的,如果它在其定义域中的每个点上都连续。

我们再给出一个等价定义。

定义 12.17 函数 $f: X \rightarrow Y$ 是连续的,如果对每个开集 $B \subset Y$, $f^{-1}(B) := \{x \in X: f(x) \in B\}$ 是开集。

当函数的值域是实数空间的子集时,这种函数被称为实值函数。就实值函数而言,还有一个连续性定义也很有用。

定义 12.18 给定函数 $f: X \rightarrow \mathbb{R}$, 对每个 $\alpha \in \mathbb{R}$, $U_\alpha = \{x \in X: f(x) \geq \alpha\}$ 为其上半集(upper contour set), $L_\alpha = \{x \in X: f(x) \leq \alpha\}$ 为其下半集(lower contour set)。

我们可以用上(下)半集表述函数的连续性。

命题 12.7 函数 $f: X \rightarrow \mathbb{R}$ 是连续的,当且仅当所有的上半集和下半集都是闭集。

极值、解和不动点*

下面的结论给出了优化问题有解的充分条件。

定理 12.1 如果 $f: X \rightarrow \mathbb{R}$ 是连续的且 X 是紧的和非空的,则存在点 $x^* = \arg \max_{x \in X} \{f(x)\}$ 。

还有另一个重要的结论。

定理 12.2 (Bolzano 介值定理) 如果 $f: [a, b] \rightarrow \mathbb{R}$ 是连续的且 $f(a) < y < f(b)$ [或 $f(b) < y < f(a)$], 则存在 $c \in (a, b)$ 使得 $f(c) = y$ 。

证明: 我们来看 y 的下半集, $L_y = \{x \in [a, b]: f(x) \leq y\}$. L_y 是非空集, 因为 $a \in L_y$. 这个集合是有界集, 因此根据完备性公理, 它有最小上界——我们将其记为点 c . 于是, $c \in L_y$ 或者 c 是个聚点. 如果 c 是个聚点, 则在 L_y 中存在某个实数列 $\{x_n\}$ 使得 $\lim x_n = c$. f 是连续函数, 因此, $\{f(x_n)\}$ 收敛于 $f(c)$. 由于对每个 n , $f(x_n) \leq y$, 所以, $f(c) \leq y$. 因此, $c \in L_y$. 再来看 y 的上半集, U_y . 由于 $b \in U_y$, U_y 是非空集. 它也是有界集, 因此, 根据完备性公理, 它有最大下界——我们将其记为点 c . 于是, $c \in U_y$ [就是说, $f(c) \geq y$] 或者 c 是个聚点. 如果 c 是个聚点, 则在 U_y 中存在某个实数列 $\{x_n\}$ 使得 $\lim x_n = c$. 由于 f 是连续函数, 因此, $\{f(x_n)\}$ 收敛于 $f(c)$. 由于对每个 n , $f(x_n) > y$, 所以, $f(c) \geq y$. 于是, $c \in U_y$. 因此, $c \in U_y \cap L_y$, 也就是说, $f(c) = y$. $\square$

在证明某些类型的方程有解时, 这一定理很有用. 事实上, 在介值定理和博弈论之间的一个自然联系便是不动点定理——不动点定理一般被用于证明均衡的存在。

定义 12.19 给定函数 $f: X \rightarrow X$, 如果点 $x \in X$ 使得 $f(x) = x$, 则 x 被称为不动点。

与此有关的一个重要结论是 Brouwer 不动点定理。

定理 12.3 (Brouwer) 如果 $X \subset \mathbb{R}^n$ 是紧的、凸的和非空的, 且 $f: X \rightarrow X$ 是连续的, 则存在不动点。

对 $n > 1$ 情况中这一定理的证明, 超出了本书的范围, 但是, 我们可以用介值定理证明一维空间中的不动点定理。

命题 12.8 如果 $f: [a, b] \rightarrow [a, b]$ 是连续的, 则它有不动点。

证明: 定义函数 $g(x) = f(x) - x$. 这是个从 $[a, b]$ 到 $[a-b, b-a]$ 的连续函数. 如果对某个 $a', b' \in [a, b]$, 有 $g(a') > 0$ 和 $g(b') < 0$ 或者 $g(a') < 0$ 和 $g(b') > 0$, 则根据介值定理, 存在某个点 $c \in [a', b'] \subset [a, b]$ 使得 $g(c) = 0$, 于是, $f(c) = c$, c 是不动点. 另有一种可能是, 对所有的 $x \in [a, b]$, $g(x) > 0$ 或者对所有的 $x \in [a, b]$, $g(x) < 0$. 也就是说, 对所有的 x , $f(x) > x$ 或对所有的 x , $f(x) < x$. 但是, 由于 $b = \sup\{x \in [a, b]\} = \sup\{f(x): x \in [a, b]\}$, $a = \inf\{x \in [a, b]\} = \inf\{f(x): x \in [a, b]\}$, 所以, 不可能存在这种情况. $\square$

在 X 是个域时 (实数集是个域), 下面的概念也很有用。

定义 12.20 设 X 是凸集, 则函数 $f: X \rightarrow \mathbb{R}$ 是拟凹的, 如果它的上半集是凸集. 就是说, 对每个 $t \in \mathbb{R}$ 和 $x, x' \in X$, 对每个 $\lambda \in (0, 1)$, 如果 $f(x) \geq t$ 且 $f(x') \geq t$, 则 $f(\lambda x + (1-\lambda)x') \geq t$. 如果最后一个不等式始终严格成立, 则函数 f 严格拟凹。

我们很容易证明拟凹目标函数的一个有用的性质。

定理 12.4 如果 X 是凸的且 $f: X \rightarrow \mathbb{R}$ 严格拟凹, 则 $\arg \max_{x \in X} \{f(x)\}$ 包含至多一个点。

12.6 对应**

对应有许多重要性质。

定义 12.21 对应 $f: X \rightarrow Y$ 是凸值的, 如果对每个 $x \in X$, 集合 $f(x)$ 是凸的。

连续性概念也可以推广至对应上。我们首先定义上像(upper image)和下像(lower image)。

定义 12.22 $E \subset Y$ 在 f 下的上像[写为 $f^+(E)$], 为集合 $f^+(E) = \{x \in X: f(x) \subset E\}$ 。

集合 E 的上像是 X 中被映射到 E 的子集上的点构成的集合。

定义 12.23 $E \subset Y$ 在 f 下的下像[写为 $f^-(E)$], 被定义为 $f^-(E) = \{x \in X: f(x) \cap E \neq \emptyset\}$ 。

集合 E 的下像是 X 中被映射到与 E 相交的集合上的点构成的集合。正如函数的连续性与 contour 集的性质有关一样, 对应的连续性与这些像集的性质有关。

定义 12.24 对应 $f: X \rightarrow Y$ 是上半连续的, 如果对每个 $x \in X$, 只要 $x \in f^+(E)$, 这里的 E 是 Y 中的开集, 就存在开球 $B(x, \varepsilon)$ 且 $B(x, \varepsilon) \subset f^+(E)$ 。

定义 12.25 对应 $f: X \rightarrow Y$ 是下半连续的, 如果对每个 $x \in X$, 只要 $x \in f^-(E)$, 这里的 E 是 Y 中的开集, 就存在开球 $B(x, \varepsilon)$ 且 $B(x, \varepsilon) \subset f^-(E)$ 。

定义 12.26 对应 $f: X \rightarrow Y$ 是连续的, 如果它是上半连续的和下半连续的。

对政治博弈论中的大部分问题, 我们可以应用比上半连续性更直观的另一个条件。

定义 12.27 对应 $f: X \rightarrow Y$ 在点 $x \in X$ 是闭的, 如果 $x_n \rightarrow x$, $y_n \in f(x_n)$ 和 $y_n \rightarrow y$ 意味着 $y \in f(x)$ 。一个对应如果在其定义域中的每个点上都是闭的, 则它是闭的。

命题 12.9 如果 Y 是紧的, 则 $f: X \rightarrow Y$ 是上半连续的当且仅当它是闭的。

下面的结论是最大值定理的一个早期形式, 它还有其他形式, 但是对正式理论而言, 其基本思想十分明确: 良好性态的优化问题的解, 对参数上的变化做出光滑反应。

定理 12.5 (Berge) 如果 $u: X \rightarrow \mathbb{R}$ 是连续函数且 $f: Y \rightarrow X$ 使得对每个 $y \in Y$, $\Gamma(y) \neq \emptyset$, 则

(1) 函数 $v: Y \rightarrow \mathbb{R}$ ——定义为 $v(y) = \max\{u(x), s. t. x \in \Gamma(y)\}$ ——是连续的; 并且(2) 对应 $a: Y \rightarrow X$ ——定义为 $a(y) = \arg \max_{x \in \Gamma(y)} \{u(x)\}$ 是上半连续的。

本节的最后一个结论是对 Brouwer 不动点定理的推广。

定理 12.6(角谷) 设 $A \subset \mathbb{R}^n$ 是紧的和凸的, $f: A \rightarrow A$ 是闭的(或上半连续的)且有非空的和凸的值, 则 f 有不动点 $x \in A$ 使得 $x \in f(x)$ 。

12.7 微积分

上述分析结论完全建立在有关拓扑性质(紧性和连续性)和凸性等知识点上。只要引入更多的结构和使用微积分做工具, 我们能够得到更精细的结论。在本节中, 我们快速浏览在本书中用到的一些微积分学基本概念。有关知识可以参考 Gill(2004)、Chiang(2003)和 Simon 与 Blume(1994)。

12.7.1 $\mathbb{R}$ 上的微分

在经验实证政治科学中, 我们所探讨的许多问题涉及在变量 x 变化时, 变量 y 如何变化的问题。如果这两个变量通过函数联系在一起, 就是说, 如果 $y = f(x)$, 则导数概念能使我们描述和量化一个变量对另一个变量的影响。假设 $y = f(x)$ 。如果把 x 增加到 $x+h$, 则有怎样的后果呢? x 每变化一个单位, 所导致的 y 的变化为:

$$\frac{\Delta y}{\Delta x} = \frac{f(x+h) - f(x)}{h}$$

这正是连接 $(x+h, f(x+h))$ 和 $(x, f(x))$ 的直线的斜率。但是, 这种衡量方法存在的问题是, 它依赖于 h 。图 12.3 中对应着 h_1 和 h_2 的两条虚线说明了这一点。我们希望有一种衡量方法, 它不依赖 h 且能够描述函数在 x 点附近的行为, 于是便有下面这个定义:

$$\frac{dy}{dx} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

即 f 对 x 的导数, 也就是图 12.3 中的实直线。我们还可以将 f 对 x 的导数写为

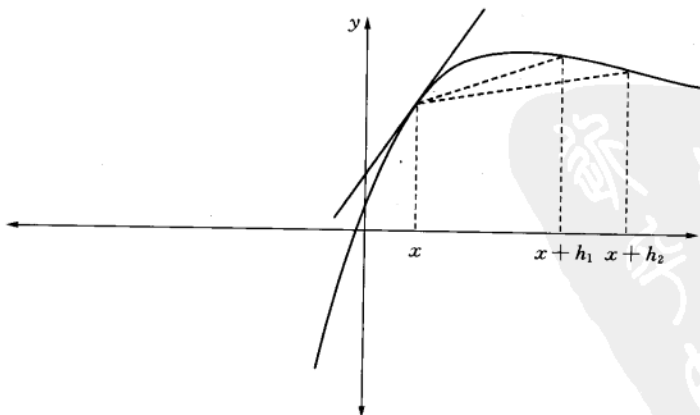


图 12.3 导数

$f'(x)$, f' 则表示 f 的导数。要注意的是, 这一极限的分子趋于零, 但是分母也趋于零, 因此它可能收敛于任何值。我们没有办法保证这一极限存在。而如果这一极限存在, 我们就说函数在 x 点可微。如果 f 在点 x 不连续, 这一极限就不可能存在。但是, f 有可能是连续的但不可微的。例如函数 $f(x) = |x|$, 它连续但是在 $x = 0$ 上是不可微的, 因为:

$$\lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h} = \lim_{h \rightarrow 0} \frac{|h|}{h}$$

这一极限并不存在, 因为数列 $h < 0$ 收敛于 -1 而数列 $h > 0$ 收敛于 1 。

导数衡量 x 变化时 y 的变化率, 所以我们可以用导数判断函数是递增函数还是递减函数。如果 $f'(x) > 0$, 则函数递增; 如果 $f'(x) < 0$, 则函数递减。

12.7.1.1 一些特殊导数

许多函数的导数有相当常见的形式, 这里给出一些函数的导数公式:

- (1) 常函数: 如果 $f(x) = c$, 则 $f'(x) = 0$ 。
- (2) 线性函数: 如果 $f(x) = a_0 + a_1x$, 则 $f'(x) = a_1$ 。
- (3) 幂函数: 如果 $f(x) = ax^n$, 则 $f'(x) = nax^{n-1}$ 。
- (4) 指数函数: 如果 $f(x) = e^{ax}$, 则 $f'(x) = ae^{ax}$ 。
- (5) 自然对数函数: 如果 $f(x) = a \ln bx$, 则 $f'(x) = \frac{a}{x}$ 。

12.7.1.2 复合函数的导数

对更复杂的函数, 如果我们能够将其分解为复合函数, 如 $f(x)$ 和 $g(x)$ 的复合函数, 就容易对其求导。下面的求导规则便利这种函数求导。

(1) 和的求导法则和差的求导法则:

$$\frac{d(f+g)}{dx} = f' + g', \quad \frac{d(f-g)}{dx} = f' - g'$$

(2) 积的求导法则:

$$\frac{d(f \times g)}{dx} = f' \times g + f \times g'$$

(3) 商的求导法则:

$$\frac{d\left(\frac{f}{g}\right)}{dx} = \frac{f' \times g - f \times g'}{g^2}$$

(4) 连续求导法则: 设 $z = g(y)$, $y = f(x)$ 且 $z = g(f(x))$, 则:

$$\frac{dz}{dx} = \frac{dz}{dy} \frac{dy}{dx} = f'(x)g'(y)$$

12.7.1.3 高阶导数

$f(x)$ 的导数(在存在时)是 x 的函数,所以,我们可以对导数求导,借此更深入地了解函数的性质。我们把 $f'(x)$ 的导数,或者说 $f(x)$ 的二阶导数,表示为 $\frac{d^2 f}{dx^2} = f''(x)$ 。和前面一样,如果 $f'' > 0$, f' 递增;如果 $f'' < 0$, f' 递减。二阶导数还能告诉我们关于函数 $f(x)$ 的特征。

- 如果 $f' > 0$ 且 $f'' > 0$,则 $f(x)$ 以递增的速度上升;
- 如果 $f' > 0$ 且 $f'' < 0$,则 $f(x)$ 以递减的速度上升;
- 如果 $f' < 0$ 且 $f'' > 0$,则 $f(x)$ 以递减的速度下降;
- 如果 $f' < 0$ 且 $f'' < 0$,则 $f(x)$ 以递增的速度下降。

图 12.4 中的函数在 x 的不同区间中,分别表现出了上述这些特征。

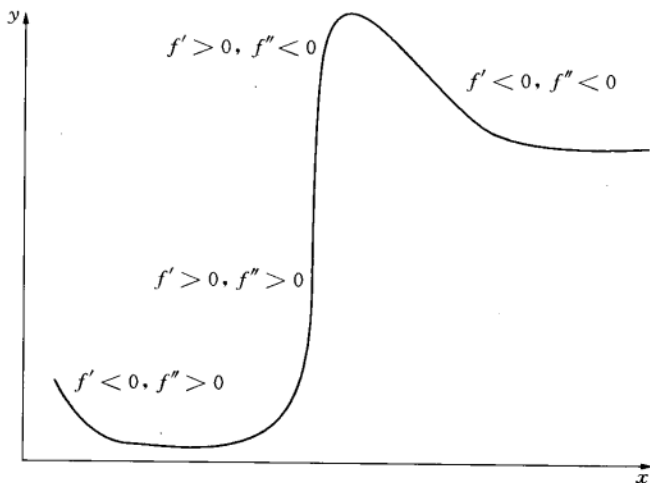


图 12.4 二阶导数

理论上说,我们可以求出第 n 阶导数,只要它存在。我们把 n 阶导数写为 $\frac{d^n f}{dx^n} = f^{(n)}(x)$ 。

12.7.1.4 函数的最大值和最小值

政治博弈论中的很多数学分析涉及最大化或最小化某些函数。选民最大化效用函数,政治家最大化选票数,政府最小化战争中阵亡人数。在找到函数的局部极值(而不是全局极值)上,导数是相当有用的工具。

直观地说,局部最大值(局部最小值)是函数停止递增(递减)而开始递减(递增)的点。因此,除非这个局部极值点同时是全局极值点并位于定义域的边界上,否则,在这个点上,导数肯定等于零。图 12.5 说明了全局最大值与局部最大值和全局最小值与局部最小值之间的区别,说明了在局部极值点上导数为什么等于零。

但是,在不是最大值或最小值的点上,导数也可能等于零,图 12.6 中包含有“鞍点”的函数说明了这一点。所以,要保证满足一阶条件的点是极值点,还要求二阶条件也得到满足。趋近局部最大值点,导数是正的;在到达这个点之后,导数为负的。就是说,导数必须递减,这意味着二阶导数不可能是正的。相反,在局部最小值点上,二阶导数不可能是负的。在鞍点上,二阶导数等于 0。

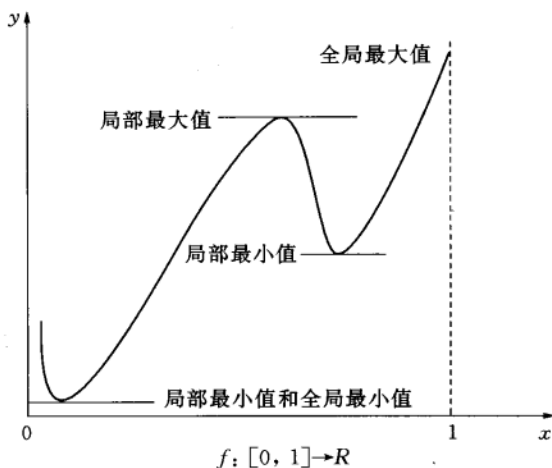


图 12.5 极值点

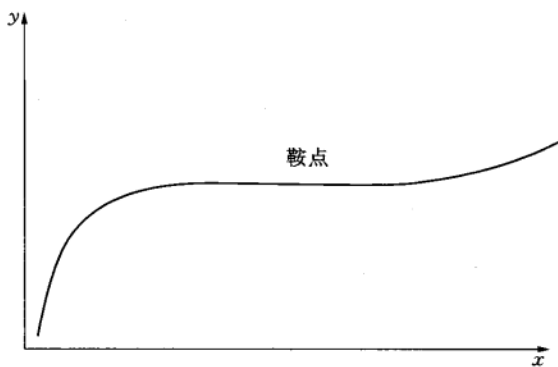


图 12.6 鞍点

一些正式定义

这里给出一些有用的定义。设 $f: D \rightarrow \mathbb{R}$, 则:

- $f(x^*)$ 是全局最大值, 如果对所有的 $x \in D$, $f(x^*) \geq f(x)$ 。
- $f(x^*)$ 是全局最小值, 如果对所有的 $x \in D$, $f(x^*) \leq f(x)$ 。
- $f(x^*)$ 是局部最大值, 如果对所有的 $\varepsilon > 0$, $|x^* - x| < \varepsilon$ 意味着 $f(x^*) \geq f(x)$ 。
- $f(x^*)$ 是局部最小值, 如果对所有的 $\varepsilon > 0$, $|x^* - x| < \varepsilon$ 意味着 $f(x^*) \leq f(x)$ 。

- 如果 $f(x^*)$ 是最大值, 则 x^* 被定义为 $\arg \max_D f(x)$ 。
- 如果 $f(x^*)$ 是最小值, 则 x^* 被定义为 $\arg \min_D f(x)$ 。
- 如果 $f'(x^*) = 0$, 则 x^* 被称为 f 的驻点。

例题 12.7 应用: 行政资源配置

某个官员有数额为 B 的预算, 计划支出于能够提高所在机构产出的两项活动上。这个机构的产出为 $O = \sqrt{x_1 x_2}$, 其中 x_1 和 x_2 分别是活动 1 和活动 2 上的支出水平。由于 $B = x_1 + x_2$, 所以可以用 $B - x_1$ 替换 x_2 , 于是, $O = \sqrt{x_1 (B - x_1)}$ 。我们现在想找到能最大化这个机构的产出的支出水平 x_1 。首先, 求出驻点值 x_1^* , 借此寻找局部最大值。产出函数的导数是:

$$O' = \frac{\left(\frac{1}{2}B - x_1^*\right)}{\sqrt{x_1^* (B - x_1^*)}}$$

令 $O' = 0$, 我们得到唯一的驻点值 $x_1^* = \frac{B}{2}$ 。为确定它是否是最大值点, 只需再求出二阶导数, 并计算它在 x_1^* 上的值。二阶导数为:

$$O'' = - \left((x_1^* (B - x_1^*))^{-\frac{1}{2}} - \left(\frac{B}{2} - x_1^*\right)^2 (x_1^* (B - x_1^*))^{-\frac{3}{2}} \right)$$

在 $x_1^* = \frac{B}{2}$ 时, 二阶导数的值为 $O'' = -2 < 0$ 。所以, $x_1^* = \frac{B}{2}$ 是局部最小值点, 实现的产出为 $\frac{B}{2}$ 。由于 $O(x_1^*) = \frac{B}{2}$ 大于 $O(0) = O(B) = 0$, 所以, 它也是全局最大值点。

12.7.1.5 函数的凹性和凸性

函数的另两个重要的性质是凹性和凸性。为说明这两个概念, 我们看图 12.7。向下弯曲的函数 f_1 , 被称为凹函数。给定任意两个点 x_1 和 x_2 , $(x_1, f(x_1))$ 和 $(x_2, f(x_2))$ 之间的连线如果位于函数下方, 这个函数是凹的。正式地说, $f: D \rightarrow \mathbb{R}$ 在集合 D 上是凹的, 当且仅当对所有的 $\lambda \in [0, 1]$ 和 $x_1, x_2 \in D$,

$$f(\lambda x_1 + (1 - \lambda)x_2) \geq \lambda f(x_1) + (1 - \lambda)f(x_2)$$

而向上弯曲的函数 f_2 , 被称为凸函数。给定任意两个点 x_1 和 x_2 , 在这两个点之间, $(x_1, f(x_1))$ 和 $(x_2, f(x_2))$ 之间的连线如果位于函数上方, 这个函数就是凸的。正式地说, $f: D \rightarrow \mathbb{R}$ 在集合 D 上是凸的, 当且仅当对所有的 $\lambda \in [0, 1]$ 和 $x_1, x_2 \in D$,

$$f(\lambda x_1 + (1 - \lambda)x_2) \leq \lambda f(x_1) + (1 - \lambda)f(x_2)$$

用严格不等号替换其中的弱不等号, 就得到了严格凹性和严格凸性的定义。

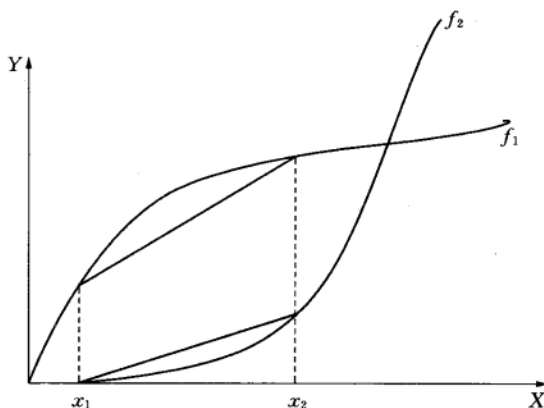


图 12.7 凸性和凹性

由于以下命题,凹函数和凸函数相当重要:

- 如果 $f: D \rightarrow \mathbb{R}$ 是凹的且 $f'(x^*) = 0$, 则 x^* 是全局最大值。
- 如果 $f: D \rightarrow \mathbb{R}$ 是凸的且 $f'(x^*) = 0$, 则 x^* 是全局最小值

这两个命题都成立, 因为, 如果 f' 存在, 则根据凹性, $f'' < 0$; 根据凸性, $f'' > 0$ 。

12.7.1.6 积分

设 $F(x)$ 是函数且满足 $F'(x) = f(x)$, 则 F 是 $f(x)$ 的反导数。我们一般把反导数表示为不定积分形式:

$$F(x) = \int f(x) dx$$

根据微分法则, 我们有以下结论 (其中, C 是任意常数)。把每个式子左边求导, 可以验证各个结论的正确性。

$$(1) \int af(x) dx = a \int f(x) dx$$

$$(2) \int (f(x) + g(x)) dx = \int f(x) dx + \int g(x) dx$$

$$(3) \int x^n dx = \frac{x^{n+1}}{n+1} + C$$

$$(4) \int \frac{1}{x} dx = \ln x + C$$

$$(5) \int e^x dx = e^x + C$$

$$(6) \int e^{f(x)} f'(x) dx = e^{f(x)} + C$$

$$(7) \int (f(x))^n f'(x) dx = \frac{f(x)^{n+1}}{n+1} + C$$

$$(8) \int \frac{f'(x)}{f(x)} dx = \ln f(x) + C$$

积分的最常见的用途是衡量函数下方区域的面积。如果 F 是 f 的反导数, 则在点 a 和点 b 之间, f 下方的面积, 为定积分:

$$\int_a^b f(x) dx = F(b) - F(a)$$

12.7.1.7 定积分求导

定积分的求导法则包括:

$$(1) \frac{d}{dx} \int_a^b f(x) dx = \int_a^b f'(x) dx$$

$$(2) \frac{d}{db} \int_a^b f(x) dx = f(b)$$

$$(3) \frac{d}{da} \int_a^b f(x) dx = -f(a)$$

$$(4) \frac{d}{d\alpha} \int_{a(\alpha)}^{b(\alpha)} f(x(\alpha)) dx = \int_a^b f'(x(\alpha)) \frac{\partial x}{\partial \alpha} dx + f(b(\alpha)) \frac{\partial b}{\partial \alpha} - f(a(\alpha)) \frac{\partial a}{\partial \alpha}$$

最后一个法则, 是最具一般性的求导法则, 被称为莱布尼茨 (Leibnitz) 法则, 是用微积分的创始人之一 Gottfried Wilhelm Leibnitz 的名字命名的。

12.7.2 多变量微分*

本节假定读者对矩阵符号有一些基本了解。矩阵是数字的长方形排列, 如:

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix}$$

对任意一个矩阵元素, 我们常用它所在的行和列表示 (例如, a_{23})。列向量是一列数字, 如:

$$B = \begin{bmatrix} b_1 \\ b_2 \end{bmatrix}$$

行向量则为一行数字, 如:

$$C = [c_1 \quad c_2 \quad c_3]$$

向量或矩阵的转置由符号“ $'$ ”注明。通过互换向量或矩阵的行和列便可得到其转置。例如:

$$A' = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix}$$

$$\mathbf{B}' = [\mathbf{b}_1 \quad \mathbf{b}_2]$$

在本节中,我们用粗字体表示向量,用非加粗字体表示标量。给定函数 $y = f(\mathbf{x})$ 。我们常常需要知道,给定 $\mathbf{x}$ 中某个坐标元素的变化, y 是如何变化的。我们来看 x_i 的偏效应:在 $\mathbf{x}$ 中的其他坐标元素保持不变时, x_i 的变化如何影响 y ? 这等价于分析函数在某一给定的切面上的行为。偏导数的正式定义是

$$\frac{\partial f}{\partial x_i} = \lim_{h \rightarrow 0} \frac{f(\mathbf{x} + \mathbf{h}_i) - f(\mathbf{x})}{h}$$

其中, $\mathbf{h}_i$ 是个向量,在其中的第 i 个位置上为 h 外,在其他所有位置上的值都为零。由于其他一切变量都被视为常数,所以,我们可以把求偏导数看作是求正常导数。

例题 12.8 设 $f(x_1, x_2) = \frac{x_1}{x_2}$, 则 $\frac{\partial f}{\partial x_1} = \frac{1}{x_2}$, $\frac{\partial f}{\partial x_2} = -\frac{x_1}{x_2^2}$ 。

我们常常把偏导数集中在一起,表示为向量形式:

$$\mathbf{D}_{\mathbf{x}}(f) = \left(\frac{\partial f}{\partial x_1}, \frac{\partial f}{\partial x_2}, \dots, \frac{\partial f}{\partial x_n} \right)'$$

称为梯度向量。梯度向量在 $\mathbf{x}$ 点的值描述了函数在 $\mathbf{x}$ 点附近的行为。

295

12.7.2.1 高阶偏导数和交叉导数

和处理单变量函数一样,我们可以求高阶偏导数,借此描述偏导数的特征。函数 f 对 x_i 的二阶偏导数为:

$$\frac{\partial}{\partial x_i} \left(\frac{\partial f}{\partial x_i} \right) = \frac{\partial^2 f}{\partial x_i^2}$$

我们可以以处理单变量函数的方式处理高阶偏导数。在多变量函数中,我们可能想知道在其他变量变化时,函数对某个变量的偏导数如何变化(改变 x_j 如何影响函数对 x_i 的偏导数?)。交叉偏导数的计算公式为:

$$\frac{\partial}{\partial x_j} \left(\frac{\partial f}{\partial x_i} \right) = \frac{\partial^2 f}{\partial x_j \partial x_i}$$

例题 12.9 设 $f(x_1, x_2) = \frac{x_1}{x_2}$ 。则 $\frac{\partial^2 f}{\partial x_1^2} = 0$, $\frac{\partial^2 f}{\partial x_2^2} = -\frac{2x_1}{x_2^3}$, $\frac{\partial^2 f}{\partial x_1 \partial x_2} = -\frac{1}{x_2^2}$,

$$\frac{\partial^2 f}{\partial x_2 \partial x_1} = -\frac{1}{x_2^2}。$$

这里要看到, $\frac{\partial^2 f}{\partial x_1 \partial x_2} = \frac{\partial^2 f}{\partial x_2 \partial x_1}$, 这一性质在一般情况中亦成立,就是说,

$$\frac{\partial^2 f}{\partial x_i \partial x_j} = \frac{\partial^2 f}{\partial x_j \partial x_i}。换句话说,偏导数的求导次序并不重要。$$

我们常把二阶偏导数和交叉偏导数集中在一起,表示为海赛矩阵的形式:

$$H(\mathbf{x}) = \begin{bmatrix} \frac{\partial^2 f(\mathbf{x})}{\partial x_1^2} & \frac{\partial^2 f(\mathbf{x})}{\partial x_1 \partial x_2} & \cdots & \frac{\partial^2 f(\mathbf{x})}{\partial x_1 \partial x_n} \\ \frac{\partial^2 f(\mathbf{x})}{\partial x_2 \partial x_1} & \frac{\partial^2 f(\mathbf{x})}{\partial x_2^2} & \cdots & \frac{\partial^2 f(\mathbf{x})}{\partial x_2 \partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f(\mathbf{x})}{\partial x_n \partial x_1} & \frac{\partial^2 f(\mathbf{x})}{\partial x_n \partial x_2} & \cdots & \frac{\partial^2 f(\mathbf{x})}{\partial x_n^2} \end{bmatrix}$$

12.7.2.2 隐函数定理*

许多均衡分析要求对参数 $y \in \mathbb{R}^k$ 的某个值,找到作为方程组 $f(\mathbf{x}; y) = 0$ 的解的 $\mathbf{x} \in \mathbb{R}^n$ 的值。在存在闭式解 $\mathbf{x}^*$ 时,我们就能求出如下形式的明确关系:

$$\mathbf{x}^* = g(y)$$

如果 g 是可微函数,则比较静态分析(即分析 y 的变化如何影响 $\mathbf{x}$)可谓直截了当。有时,我们能够证明对每个 y ,存在解 $\mathbf{x}^*$,但是,我们无法直接得到函数 $g(\cdot)$ 。例如,不动点定理可能告诉我们方程组 $f(\mathbf{x}; y) = 0$ 的解存在,但是我们可能无法以解析的方式求出它的解,即把向量 $\mathbf{x}$ 表示为 y 的函数。

在恰当条件下,隐函数定理能让我们隐含地求出导数 $D_y \mathbf{x}^*$ 。我们首先在只有一个内生变量和一个外生变量的情况下表述这一定理。

定理 12.7 (隐函数定理) 给定 $y \in \mathbb{R}$, 设 $\mathbf{x}^* \in \mathbb{R}$ 是 $f(\mathbf{x}, y) = 0$ 的解。如果 $f(\cdot, \cdot)$ 连续可微且 $\frac{\partial f(\mathbf{x}^*, y^*)}{\partial x} \neq 0$, 则对包含 $\mathbf{x}^*$ 的某个开集 A 和包含 y^* 的某个开集 B , 存在连续可微函数 $\phi: B \rightarrow A$ 使得 $f(\phi(y), y) = 0$ 。这一函数在 y^* 上的导数为

$$\frac{\partial \phi(y^*)}{\partial y} = - \frac{\frac{\partial f(\mathbf{x}^*, y^*)}{\partial y}}{\frac{\partial f(\mathbf{x}^*, y^*)}{\partial x}}$$

为得到更一般的隐函数定理,设内生向量为 $\mathbf{x} = (x_1, \dots, x_n) \in \mathbb{R}^n$, 外生向量为 $\mathbf{y} = (y_1, \dots, y_k) \in \mathbb{R}^k$ 。设方程组 $f(\mathbf{x}, \mathbf{y}) = 0$ 的形式为:

$$\begin{aligned} f_1(x_1, \dots, x_n; y_1, \dots, y_k) &= 0 \\ &\vdots \\ f_n(x_1, \dots, x_n; y_1, \dots, y_k) &= 0 \end{aligned}$$

这一方程组在内生变量上的 $n \times n$ 雅可比矩阵,为梯度向量的转置所构成的 $n \times n$ 矩阵:

$$J = D_{\mathbf{x}} f(\mathbf{x}^*, \mathbf{y}^*) = \begin{bmatrix} D_{\mathbf{x}} f'_1 \\ \vdots \\ D_{\mathbf{x}} f'_n \end{bmatrix}$$

在矩阵运算中,与除运算相对应的是所谓的“取逆”运算。对任意的 $n \times n$ 矩阵 M , 矩阵 M^{-1} 满足方程 $MM^{-1} = I$, 其中 I 是单位矩阵,在其主对角线上所有元素为 1, 在其他位置上所有元素为 0。一个矩阵,如果它的逆存在,则称之为非奇异型矩阵。关于逆矩阵的计算的详细介绍,建议读者参考 Chiang(2004)以及 Simon 和 Blume(1994)或任何一本初级线性代数课本。

$n \times k$ 矩阵

$$D_y f(x^*, y^*) = \begin{bmatrix} D_y f_1' \\ \vdots \\ D_y f_n' \end{bmatrix}$$

是由 $f(\cdot, \cdot)$ 中所有式子对外生变量求导而得到的导数的转置。

定理 12.8 (隐函数定理) 设 x^* 是 $f(x, y) = 0$ 在 $y \in \mathbb{R}^k$ 上的解。如果 $f_1(\cdot)$ 到 $f_n(\cdot)$ 在 x 和 y 的每个坐标上都连续可导,并且这个方程组对内生变量的雅克比矩阵是非奇异的,则对包含 x^* 的某个开集 A 和包含 y^* 的某个开集 B , 存在连续可微函数 $\phi: B \rightarrow A$ 使得 $f(\phi(y), y) = 0$ 。这一函数在 y^* 上的导数为 $n \times k$ 矩阵:

$$D_y \phi(y^*) = -[D_x f(x^*, y^*)]^{-1} D_y f(x^*, y^*)$$

12.7.2.3 $\mathbb{R}^n$ 中的优化

前面说过,如果要最大化函数 $f: \mathbb{R} \rightarrow \mathbb{R}$, 我们只需求解满足 $f'(x^*) = 0$ 的 x 值。^①如果这一条件不成立,则在 x^* 的某邻域内存在其他 x , $f(x)$ 的值比 $f(x^*)$ 大。

这一点对多变量函数的优化亦成立。在多变量函数的优化中,函数对 x 的每个坐标元素的导数必须为 0。设 $\frac{\partial f}{\partial x_i} > 0$, 则函数值随 x_i 的微小增长而增长,随 x_i 的微小下降而下降。同样的道理,在内点最优解上,不可能有 $\frac{\partial f}{\partial x_i} < 0$ 。所以, x^* 优化 $f: \mathbb{R}^n \rightarrow \mathbb{R}$ 的必要条件是

$$Df(x^*) = 0$$

最大值和最小值的二阶条件则建立在海赛矩阵基础上,需要矩阵代数中的一些更高深的概念。最大值(最小值)的充分条件是海赛矩阵 H 是正定的(负定的)。

定义 12.28 $n \times n$ 矩阵 M 是正定的,如果对所有向量 $v \in \mathbb{R}^n$, $v'Mv > 0$; 是负定的,如果对所有 $v \in \mathbb{R}^n$, $v'Mv < 0$ 。这两个不等式如果都为弱不等式,则这个矩阵就是正半定的和负半定的。

我们首先对 $\mathbb{R}^2$ 给出以下结论。

定理 12.9 $x^* \in \mathbb{R}^2$ 是局部最大值元素,如果 $Df(x^*) = 0$, $\frac{\partial^2 f}{\partial x_1^2} < 0$, $\frac{\partial^2 f}{\partial x_2^2} < 0$,

^① 由于这里的定义域是 $\mathbb{R}$, 所以不存在角解的问题。

$$\text{且 } \frac{\partial^2 f}{\partial x_1^2} \frac{\partial^2 f}{\partial x_2^2} > \left(\frac{\partial^2 f}{\partial x_1 \partial x_2} \right)^2.$$

定理 12.10 $\mathbf{x}^* \in \mathbb{R}^2$ 是局部最小值元素, 如果 $Df(\mathbf{x}^*) = 0$, $\frac{\partial^2 f}{\partial x_1^2} > 0$,

$$\frac{\partial^2 f}{\partial x_2^2} > 0, \text{ 且 } \frac{\partial^2 f}{\partial x_1^2} \frac{\partial^2 f}{\partial x_2^2} > \left(\frac{\partial^2 f}{\partial x_1 \partial x_2} \right)^2.$$

高维空间中的结论与此类似。

定理 12.11 $\mathbf{x}^* \in \mathbb{R}^2$ 是局部最大值元素, 如果 $Df(\mathbf{x}^*) = 0$ 且海赛矩阵 $H(\mathbf{x}^*)$ 是负定的。

定理 12.12 $\mathbf{x}^* \in \mathbb{R}^2$ 是局部最小值元素, 如果 $Df(\mathbf{x}^*) = 0$ 且海赛矩阵 $H(\mathbf{x}^*)$ 是正定的。

如果 $Df(\mathbf{x}^*) = 0$ 且海赛矩阵 $H(\mathbf{x}^*)$ 是负半定的, 则 $\mathbf{x}^*$ 可能是局部最大值元素。但这一点未必始终成立。对正半定海赛矩阵亦有类似结论。但是, 下面的必要性条件肯定成立。

定理 12.13 如果 $\mathbf{x}^* \in \mathbb{R}^n$ 是局部最大值元素且目标函数二阶可导, 则海赛矩阵 $H(\mathbf{x}^*)$ 是负半定的。

定理 12.14 如果 $\mathbf{x}^* \in \mathbb{R}^n$ 是局部最小值元素且目标函数二阶可导, 则海赛矩阵 $H(\mathbf{x}^*)$ 是正半定的。

例题 12.10 党派内政治资源的配置

设某个政党想把一笔资金分配于两次议会选举活动中。此党派认为两次选举中所争得的议会席位的价值分别是 W_1 和 W_2 (从所失去的每个议会席位中, 这个党派得到的收益是 0)。设 $\frac{x_i}{1+x_i}$ 是这个党派赢得席位 i 的概率, 其中 x_i 是在选举活动 i 上支出的资金数量。竞选支出 x_i 的成本就是 x_i 。这个党派想通过对 (x_1, x_2) 的选择, 最大化:

$$\frac{x_1}{1+x_1} W_1 + \frac{x_2}{1+x_2} W_2 - x_1 - x_2$$

一阶条件为:

$$\frac{W_1}{(1+x_1)^2} - 1 = 0$$

$$\frac{W_2}{(1+x_2)^2} - 1 = 0$$

二阶导数的海赛矩阵为:

$$\begin{bmatrix} -\frac{W_1}{(1+x_1)^3} & 0 \\ 0 & -\frac{W_2}{(1+x_2)^3} \end{bmatrix}$$

从一阶条件中, 得到四个可能的驻点:

$$(-\sqrt{W_1}-1, -\sqrt{W_2}-1), (\sqrt{W_1}-1, -\sqrt{W_2}-1), \\ (-\sqrt{W_1}-1, \sqrt{W_2}-1) \text{ 和 } (\sqrt{W_1}-1, \sqrt{W_2}-1)$$

二阶条件则要求 $-\frac{W_i}{(1+x_i)^3} < 0$, 即 $x_i > -1$ 。由于 $W_i > 0$, 所以, 满足二阶条件的唯一的驻点为 $(\sqrt{W_1}-1, \sqrt{W_2}-1)$ 。

12.7.2.4 凹函数和凸函数

函数的凹性和凸性定义很容易推广到 $\mathbb{R}^n$ 中。

定义 12.29 设 $U \subseteq \mathbb{R}^n$ 且 $f: U \rightarrow \mathbb{R}$ 。则函数 f 是凹的, 如果对所有的 $x, y \in U$ 和 $\lambda \in [0, 1]$, $f(\lambda x + (1-\lambda)y) \geq \lambda f(x) + (1-\lambda)f(y)$ 。

定义 12.30 设 $U \subseteq \mathbb{R}^n$ 且 $f: U \rightarrow \mathbb{R}$ 。则函数 f 是凸的, 如果对所有的 $x, y \in U$ 和 $\lambda \in [0, 1]$, $f(\lambda x + (1-\lambda)y) \leq \lambda f(x) + (1-\lambda)f(y)$ 。

和前面一样, 凹性确保驻点产生全局最大值, 凸性则保证了全局最小值。

定理 12.15 设 $f: U \rightarrow \mathbb{R}$ 是二阶可微函数, 其中 U 是 $\mathbb{R}^n$ 的开的和凸的子集。如果 f 是 U 上的凹函数且对 $x^* \in U$, $Df(x^*) = 0$, 则 x^* 是 f 在 U 上的全局最大值元素。如果 f 是 U 上的凸函数且对 $x^* \in U$, $Df(x^*) = 0$, 则 x^* 是 f 在 U 上的全局最小值元素。

12.7.2.5 约束条件下的最大化

1. 等式约束条件

在政治博弈论的许多背景中, 我们需要求解约束条件下的最大化问题。这种约束可能是决策者选择集上的可行性约束条件, 也可能源于其他人的行为。这种最大化问题的形式为:

$$\begin{aligned} \max f(x) \\ \text{s. t. } g_1(x) = 0 \\ g_2(x) = 0 \\ \vdots \\ g_k(x) = 0 \end{aligned}$$

其中, 函数 $f: \mathbb{R}^n \rightarrow \mathbb{R}$ 和各个函数 $g_j: \mathbb{R}^n \rightarrow \mathbb{R}$ 都是二阶可微的。通过建立和求解无约束条件的优化问题, 可以找到这种约束优化问题的解。这里的技巧是使约束条件进入到目标函数。

拉格朗日函数:

$$L(x, \lambda) = f(x) - \sum_{j=1}^k \lambda_j g_j(x)$$

便是变换后得到的目标函数。它的大小决定于原始问题中的选择变量 x 和由 k 个变量构成的一个新向量。这些新的变量是约束条件的乘数(每个约束条件都有自

己的乘数)。原始问题中有一个实值目标函数和 k 个约束条件,而在转换后的最大化问题中,目标函数由 $k+1$ 个实值函数的和构成。拉格朗日函数的优化,需满足的一阶条件为:

$$\frac{\partial f(\mathbf{x})}{\partial x_i} = \sum_{j=1}^k \lambda_j \frac{\partial g_j(\mathbf{x})}{\partial x_i}, i = 1, \dots, n$$

$$g_j(\mathbf{x}) = 0, j = 1, \dots, k$$

其中的前 n 个条件给出了这一约束优化问题的解 $\mathbf{x}$ 需满足的必要条件。更正式地说:

定理 12.16 (拉格朗日定理) 假设 k 个约束函数的梯度向量是线性无关向量。如果 $\mathbf{x}^*$ 是这一约束条件下的最大化问题的解,则存在拉格朗日乘数向量 $\lambda \in \mathbb{R}^k$, 使得 $(\mathbf{x}^*, \lambda)$ 是上述一阶条件的解。

要理解把约束条件下的最大化问题转化为无约束条件的最大化问题背后的原理,最好的办法是分析拉格朗日函数最大化中的前 n 个一阶条件。这些条件要求, (在一阶条件的解上) 改变 $\mathbf{x}$ 而导致的 f 值的任何增加, 必须导致至少一个约束函数 g 的值上的相应变化。换句话说, 如果 $\mathbf{x}$ 是拉格朗日最大化问题的解, 则 f 值上的任何增加都必然违反约束条件。这种求解方法要起作用, 还必须满足其中的线性无关条件, 即约束条件对变量 $\mathbf{x}$ 的雅可比矩阵的秩为 k (就是说, k 个约束条件的梯度向量是线性无关的)。没有这一约束限制条件, $\sum_{j=1}^k \lambda_j \frac{\partial g_j(\mathbf{x})}{\partial x_i}$ 上的变化未必意味着约束条件被违背。

例题 12.11 再谈党派资源配置

假设前述政党有预算约束。在把资金跨选区分配时, 它必须满足预算约束。就是说, 该党派现在要最大化:

$$\frac{x_1}{1+x_1} W_1 + \frac{x_2}{1+x_2} W_2$$

约束条件是 $B = x_1 + x_2$ 。

拉格朗日函数是:

$$\frac{x_1}{1+x_1} W_1 + \frac{x_2}{1+x_2} W_2 + \lambda(x_1 + x_2 - B)$$

一阶条件为:

$$\frac{W_1}{(1+x_1)^2} - \lambda = 0$$

$$\frac{W_2}{(1+x_2)^2} - \lambda = 0$$

$$x_1 + x_2 = B$$

从前两个条件中,得到:

$$\frac{W_2}{W_1} = \frac{(1+x_2)^2}{(1+x_1)^2}$$

即

$$\sqrt{\frac{W_2}{W_1}} = \frac{1+x_2}{1+x_1}$$

与预算约束结合,就有了两个方程和两个未知数。由于只有正根才有意义,所以,

$$1 + \sqrt{\frac{W_2}{W_1}} = \frac{1+B-x_1}{1+x_1}$$

于是:

$$x_1^* = \frac{1+B-\sqrt{\frac{W_2}{W_1}}}{\sqrt{\frac{W_2}{W_1}}+1}$$

和

$$x_2^* = \frac{1+\sqrt{\frac{W_2}{W_1}}(B-1)}{\sqrt{\frac{W_2}{W_1}}+1}$$

2. 不等式约束条件

前述问题要求约束条件形式为 $g_j(\mathbf{x}) = 0$ 。还有一大类优化问题只要求一组不等式约束或等式约束得到满足。这类问题的一般形式为:

$$\begin{aligned} & \max f(\mathbf{x}) \\ & \text{s. t. } g_j(\mathbf{x}) = 0, j = 1, \dots, k \\ & \quad h_t(\mathbf{x}) \leq 0, t = 1, \dots, w \end{aligned}$$

我们依然假定有关函数都是可微的。Kuhn-Tucker 条件类似拉格朗日条件,两者之间的差异在于对不等式约束条件的处理方式。转换后得到的一阶条件为:

$$\text{对每个 } i = 1, \dots, n, \quad \frac{\partial f(\mathbf{x})}{\partial x_i} = \sum_{j=1}^k \lambda_j \frac{\partial g_j(\mathbf{x})}{\partial x_i} + \sum_{t=1}^w \lambda_t \frac{\partial h_t(\mathbf{x})}{\partial x_i}$$

$$\text{对每个 } j = 1, \dots, k, \quad g_j(\mathbf{x}) = 0$$

$$\text{对每个 } t = 1, \dots, w, \quad \lambda_t h_t(\mathbf{x}) = 0$$

Kuhn-Tucker 条件和拉格朗日条件之间的区别体现在对不等式约束条件的处理上,在 Kuhn-Tucker 条件中,约束条件或者是紧的[就是说, $h_t(\mathbf{x}) = 0$] 或者乘数 λ_t 为 0。

3. 包络定理*

在应用中,目标函数或者约束条件还可能决定于外生变量 $y = (y_1, \dots, y_l, \dots, y_z) \in \mathbb{R}^z$ 。我们来看下面的问题:

$$\begin{aligned} \max & f(x; y) \\ \text{s. t. } & g_j(x; y) = 0, j = 1, \dots, k \\ & h_t(x, y) \leq 0, t = 1, \dots, w \end{aligned}$$

定义 $v(y)$ 为值函数,它是映射 $v: \mathbb{R}^z \rightarrow \mathbb{R}$, $v(y) = f(x^*(y); y)$, 其中, $x^*(y)$ 是最大化问题的解。最大值定理指出,在恰当条件下,这个值函数是连续的。我们可利用微分工具,对值函数与外生参数的依赖关系,得到更深入的认识。

定理 12.17 假设 $v(y')$ 在 y' 上可微, $(x^*(y'), \lambda(y'))$ 是上述最大化问题的解,并且在包含 $x^*(y')$ 的某个开集 A 和包含 y' 的某个开集 B 上,在解 $x^*: B \rightarrow A$ 上紧的约束条件集是固定的,则对每个 $i = 1, \dots, n$:

$$\frac{\partial v(y')}{\partial y_i} = \frac{\partial f(x^*(y'); y)}{\partial y_i} - \sum_{j=1}^k \lambda_j \frac{\partial g_j(x^*(y'); y)}{\partial y_i} - \sum_{t=1}^w \lambda_t \frac{\partial h_t(x^*(y'); y)}{\partial y_i}$$

利用这一定理,在求 $\frac{\partial v(y')}{\partial y_i}$ 时,我们不必考虑 $D_{y_i} x^*(y)$ 。

例题 12.12 再看前述党派如何配置其数量为 y 的资源到 k 个选区。假设对选区 $i = 1, \dots, k$, 它分配的资金数量分别为 $r_i \in \mathbb{R}$ 。由于它只有数量为 y 的资源可用于竞选,所以,约束条件 $y - \sum_{i=1}^k r_i \geq 0$ 必须得到满足。设 $f(r): \mathbb{R}^k \rightarrow \mathbb{R}$ 为这个党派从某种资源配置 r 中得到的收益。根据包络定理,如果 $v(y)$ 是从求解优化问题中得到的值函数, λ 是约束条件的乘数,则 $\frac{\partial v(y)}{\partial y} = \lambda$ 。就是说,乘数 λ 是这一党派资源的微小增加所实现的边际值。

12.7.2.6 多重积分

我们用多变量定积分计算多变量函数下方的面积:

$$\int_{a_1}^{b_1} \cdots \int_{a_n}^{b_n} f(x_1, \dots, x_n) dx_1 \cdots dx_n$$

在计算多变量积分时,我们依次保持其他变量不变而对其中一个变量求积分。设我们对 x_1 求积分。如果 $F_1(x_1, \dots, x_n)$ 是对 x_1 的偏反导数,则:

$$\begin{aligned} \int_{a_1}^{b_1} \cdots \int_{a_n}^{b_n} f(x_1, \dots, x_n) dx_1 \cdots dx_n &= \int_{a_2}^{b_2} \cdots \int_{a_n}^{b_n} F_1(b_1, \dots, x_n) dx_2 \cdots dx_n \\ &\quad - \int_{a_2}^{b_2} \cdots \int_{a_n}^{b_n} F_1(a_1, \dots, x_n) dx_2 \cdots dx_n \end{aligned}$$

我们持续进行这一过程,求 F_1 对 x_2 的偏反导数……在这里,先计算哪一个定偏积分并不重要。

例题 12.13 计算 $\int_{\frac{1}{2}}^2 \int_{\frac{1}{2}}^1 x^2 y dx dy$ 。首先求得 $F_1(x, y) = \frac{x^3 y}{3}$, 于是:

$$\begin{aligned} \int_{\frac{1}{2}}^2 \int_{\frac{1}{2}}^1 x^2 y dx dy &= \int_{\frac{1}{2}}^1 \frac{8}{3} y dy - \int_{\frac{1}{2}}^1 \frac{1}{3} y dy \\ &= \frac{8}{3} \left[\frac{1}{2} - \frac{1}{8} \right] - \frac{1}{3} \left[\frac{1}{2} - \frac{1}{8} \right] \\ &= \frac{7}{8} \end{aligned}$$

12.8 概率论

在第3章中我们看到,不确定性条件下的决策模型对概率论的依赖程度很高。本节介绍概率论的基本知识并证明其一些重要结论。

12.8.1 结果和事件

概率论中的最基本概念是结果和事件。设 S 是某个随机过程产生的所有可能结果构成的集合,称为**样本空间**。元素 $s \in S$ 被称为一个**结果**。

例题 12.14 抛两枚硬币: $S = \{HH, HT, TH, TT\}$ 。

例题 12.15 失业率: $S = [0, 100]$ 。

例题 12.14 中的样本空间是离散的,其中结果的数量有限;例题 12.15 中的样本空间是连续的,其中结果的数量是无限的。

给定样本空间,我们定义一个事件为一个子集 $A \subseteq S$ 。所以,事件是由一些结果构成的集合。

例题 12.16 $A = \{TH, HT\}$: “抛第二枚硬币得到的结果与抛第一枚硬币得到的结果不同”。

例题 12.17 $S = [4, 13]$: “失业率在 4% 到 13% 之间”。

12.8.2 概率论的几个公理

概率论探讨的是各种事件的发生几率。我们用 $\Pr(A)$ 表示事件 A 的发生概率。概率论建立在关于 $\Pr(A)$ 的以下公理基础上。

公理 12.1 对任意事件 A , $\Pr(A) \geq 0$ 。

公理 12.2 $\Pr(S) = 1$

公理 12.3 设 $A_1, A_2, \dots$ 是无限个不相交的事件, 则 $\Pr(\bigcup_{i=1}^{\infty} A_i) = \sum_{i=1}^{\infty} \Pr(A_i)$ 。

公理 12.1 指出, 任何事件的概率都是非负的; 公理 12.2 指出, 存在某个事件, 它的发生概率为 1。公理 12.3 谈的是相互排斥的或相互不交的事件族的概率, 它指出, 一组相互排斥的事件的概率等于其中各个事件的概率之和。这些公理直接产生了关于概率的许多有用的性质。

空事件的概率为 0。

定理 12.18 $\Pr(\emptyset) = 0$ 。

公理 12.3 可直接应用于有限数量的不交事件的发生概率。

定理 12.19 设 $A_1, A_2, \dots, A_n$ 是有限个不交事件, 则 $\Pr(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n \Pr(A_i)$ 。

这个定理, 加上公理 12.1 和公理 12.2, 意味着某个事件的概率和这个事件的补的概率之和为 1。

定理 12.20 设 $S|A$ 是 A 的补, 则 $\Pr(A) + \Pr(S|A) = 1$ 。

这个定理的一个直接推论是任何事件的概率(弱)小于 1。

定理 12.21 给定任意事件 A , $0 \leq \Pr(A) \leq 1$ 。

如果事件 B 同时是事件 A 的真子集, 则事件 A 的概率至少和事件 B 的概率一样大。

定理 12.22 如果 $B \subseteq A$, 则 $\Pr(A) \geq \Pr(B)$ 。

下面的两个定理与事件的并的概率有关。给定事件 A 和 B , $A \cup B$ 可以被分解为三个不交集 $A \setminus B$, $B \setminus A$ 和 $A \cap B$ 。根据定理 12.19, $\Pr(A \cup B) = \Pr(A \setminus B) + \Pr(B \setminus A) + \Pr(A \cap B)$ 。根据这一定理, $\Pr(A) = \Pr(A \setminus B) + \Pr(A \cap B)$, $\Pr(B) = \Pr(B \setminus A) + \Pr(A \cap B)$ 。定理 12.23 概括了这些结论。

定理 12.23 给定任意两个事件 A 和 B , $\Pr(A \cup B) = \Pr(A) + \Pr(B) - \Pr(A \cap B)$ 。

定理 12.24 是定理 12.23 的直接推论。

定理 12.24 给定任意 n 个事件 $A_1, A_2, \dots, A_n$:

$$\Pr(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n [\Pr(A_i) - \sum_{j>i}^n \Pr(A_i \cap A_j) + \sum_{k>j>i}^n \Pr(A_i \cap A_j \cap A_k) - \dots]$$

独立性和条件概率

我们现在探讨各个事件的发生概率之间的关系。我们想知道的是, 某个事件的发生如何影响另一事件的概率。给定事件为 A 和 B 。假设我们知道事件 B 已发生; 在此前提下, 事件 A 的发生概率是多少?

一种可能是, 这两个事件的概率之间没有任何关系。如果 $\Pr(A \cap B) = \Pr(A)\Pr(B)$, 我们说两个事件 A 和 B 相互独立。当两个事件相互独立时, 其中一个

事件发生的事实,对另一事件的概率没有任何影响。假设事件 B 发生, $\Pr(A \cap B)$ 等于 $\Pr(A)$ 。就是说,事件 A 发生与否,不受事件 B 发生与否影响。在这一逻辑下,我们有了下面的独立性的一般定义。

定义 12.31 设 $A_1, A_2, \dots, A_n$ 是一组事件。如果 $\Pr(\bigcap_{i=1}^n A_i) = \prod_{i=1}^n \Pr(A_i)$, 则这些事件相互独立。

这一定义要成立,这一事件族的任意一个子集也必须独立。DeGroot 和 Schervish(2001)给出了一个例子,其中,三个事件虽两两独立但是这三个事件却不相互独立。

我们来讨论各个事件之间存在依赖性的情况。分析事件之间这种关系的最重要的概念是条件概率。设有事件 A 和 B 。给定事件 B ,事件 A 的条件概率是,在 B 已经发生时事件 A 的发生概率。我们把给定事件 B ,事件 A 的条件概率表示为:

$$\Pr(A | B) = \frac{\Pr(A \cap B)}{\Pr(B)}, \text{ 其中 } \Pr(B) > 0$$

如果 A 和 B 相互独立,则 $\Pr(A | B) = \Pr(A)$ 。并且,条件概率满足:

$$\Pr(A \cap B) = \Pr(A | B)\Pr(B) = \Pr(B | A)\Pr(A)$$

12.8.3 贝叶斯法则

概率论在政治博弈论中最重要的应用之一是预测行为主体如何利用所观察到事件对未观察到的事件的发生概率做出估计。贝叶斯法则告诉我们,这种估计是如何做出的。要陈述和证明这一法则,还需要增加定义。分割是覆盖整个样本空间的相互排斥的事件构成的集合。

定义 12.32 样本空间 S 的分割,是一组不交的事件 $A_1, \dots, A_k$ 且 $\bigcup_{i=1}^k A_i = S$ 。

我们现在给出贝叶斯法则。

定理 12.25 如果 $A_1, \dots, A_k$ 是 S 的一个分割并且 $\Pr(B) > 0$, 则对任何的 $B \subset S$ 和 $A_j \subset S$,

$$\Pr(A_j | B) = \frac{\Pr(A_j)\Pr(B | A_j)}{\sum_{i=1}^k \Pr(A_i)\Pr(B | A_i)}$$

证明: 我们分几个步骤证明这个定理。

结论 1: 如果 $A_1, \dots, A_k$ 是 S 的分割,且 B 是 S 的子集,则集合 $A_1 \cap B, \dots, A_k \cap B$ 构成 B 的分割,因为 $\bigcup_{i=1}^k (A_i \cap B) = \bigcup_{i=1}^k A_i \cap B = S \cap B = B$ 。

结论 2: 如果 $A_1, \dots, A_k$ 是 S 的分割,则 $\Pr(B) = \sum_{i=1}^k \Pr(A_i \cap B)$ 。这是结论 1 和定理 12.19 的直接应用。

结论 3: 如果 $A_1, \dots, A_k$ 是 S 的分割且对所有的 i , $\Pr(A_i) > 0$, 则 $\Pr(B) = \sum_{i=1}^k \Pr(B | A_i)\Pr(A_i)$ 。这是结论 2 和条件概率的乘法法则的直接应用。

我们现在来证明上述定理。如果 $\Pr(B) > 0$, 根据条件概率的定义:

$$\Pr(A_j | B) = \frac{\Pr(A_j \cap B)}{\Pr(B)}$$

把 $\Pr(A_j \cap B) = \Pr(B | A_j)\Pr(A_j)$ 代入到分子, 把 $\Pr(B) = \sum_{i=1}^k \Pr(B | A_i)\Pr(A_i)$ 代入到分母, 我们就得到了贝叶斯法则。□

12.8.4 随机变量和分布

有时利用数字表示结果和事件更加方便。这种表示被称为随机变量。实值随机变量实际上是个函数, 它把所有可能的结果映射到实数上。

定义 12.33 给定样本空间 S , $X: S \rightarrow \mathbb{R}$ 是个随机变量, 对每个可能的结果 $s \in S$, 它赋予一个实数 $X(s)$ 。

利用随机变量的定义, 可以直接把事件定义为实数集, 并定义随机变量上的概率分布。分布不过是对这种事件赋予概率。

定义 12.34 设 A 是 $\mathbb{R}$ 的任意子集, $\Pr(X \in A)$ 表示 X 落于 A 中的概率, 于是, $\Pr(X \in A) = \Pr\{s: X(s) \in A\}$ 。 X 的概率分布对所有的 $A \subset \mathbb{R}$, 给出概率 $\Pr(X \in A)$ 。

12.8.4.1 离散分布

随机变量 X 如果它只有有限个结果 $x_1, x_2, \dots, x_k$, 有离散分布。所有可能的结果构成的集合被称为这一分布的支持。

定义 12.35 某个随机变量 X 如果有离散分布, 则 X 的概率函数被定义为 $f(x) = \Pr(X = x)$, 其中, x 为实数。

如果 x 不等于 X 的支持中的任意一点, 则 $f(x) = 0$ 。根据概率论公理, $0 \leq \sum_{i=1}^k f(x_i) \leq 1$ 且 $0 \leq f(x_i) \leq 1$ 。

例题 12.18 整数上的均匀分布

假设 X 在 k 个整数中——如 $1, 2, 3, \dots, k$ 中——等可能地取值, 则 X 的概率函数为:

$$f(x) = \begin{cases} \frac{1}{k}, & \text{如果 } x = 1, \dots, k \\ 0, & \text{在其他情况中} \end{cases}$$

例题 12.19 两项分布

假设一项实验成功的概率为 p , 失败的概率为 $1 - p$ 。则在 n 次实验中成功 x 次的概率为:

$$f(x) = \begin{cases} \binom{n}{x} p^x (1-p)^{n-x}, & \text{如果 } x = 0, \dots, n \\ 0, & \text{在其他情况中} \end{cases}$$

其中, $\binom{n}{x} = \frac{n!}{x!(n-x)!}$ 。

12.8.4.2 连续分布

X 如果在连续统上取值, 则为连续随机变量。如果 X 是连续随机变量, 则存在非负函数 f , 使得对任何区间 $A = [a, b]$:

$$\Pr(X \in A) = \int_A f(x) dx = \int_a^b f(x) dx$$

函数 f 被称为概率密度函数(pdf)。它没有告诉我们 $\Pr(X = x)$ 的值(它等于 0), 但告诉我们, $f(x) = \lim_{\varepsilon \rightarrow 0} \Pr(X \in [x - \varepsilon, x + \varepsilon])$ 。每个概率密度函数都满足以下条件:

对所有的 x , $f(x) \geq 0$

$$\int_{-\infty}^{\infty} f(x) dx = 1$$

集合 $X^s = \{x: f(x) > 0\}$ 被称为 X 的支持。

例题 12.20 区间上的均匀分布

设 a 和 b 是两个实数。实验内容是从 $S = [a, b]$ 中选取点 X , 设 X 属于任意一个子区间的概率与它的长度成正比。就是说, 在 S 中的任何一点上, 概率密度肯定都相同, 而在区间 S 以外的任何点上, 概率密度为 0。因此, $\int_{-\infty}^{\infty} f(x) dx = \int_a^b f(x) dx = \int_a^b c dx = 1$ 。求解 c , 得到 $c = \frac{1}{b-a}$ 。就是说, 这个均匀分布的概率密度函数为:

$$f(x) = \begin{cases} \frac{1}{b-a}, & \text{如果 } x \in [a, b] \\ 0, & \text{在其他情况中} \end{cases}$$

12.8.4.3 累积分布函数

累积分布函数(cdf)是个实值函数, 对任意实数 x , 它给出了 X 取值不大于 x 的概率:

$$F(x) = \Pr(X \leq x)$$

对离散分布, $F(x) = \sum_{\{i: x_i \leq x\}} f(x_i)$ 。对连续分布, $F(x) = \int_{-\infty}^x f(\xi) d\xi$ 。在任一点上, 只要 F 可微, 就有 $F'(x) = f(x)$ 。因此, 有密度函数的连续随机变量既可以用概率密度函数也可以用累积分布函数表示。

以下是累积分布函数的一些重要特征:

(1) $\Pr(X > x) = 1 - F(x)$ 。

$$(2) \Pr(x_2 > X \geq x_1) = F(x_2) - F(x_1).$$

(3) F 是非递减的(就是说,如果 $x_2 > x_1$, 则 $F(x_2) \geq F(x_1)$).

$$(4) \lim_{x \rightarrow -\infty} F(x) = 0, \lim_{x \rightarrow \infty} F(x) = 1.$$

(5) F 总是右连续。如果点 x 以正概率发生,则 F 在点 x 上可能不是左连续的。见图 12.8。

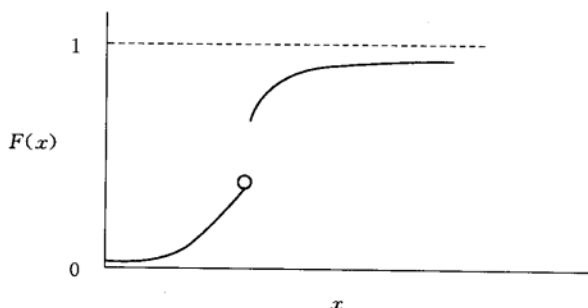


图 12.8 不连续的累积分布函数

12.8.4.4 二元分布

我们有时候要知道两个或多个随机变量,如 X 和 Y ,同时取某些值的概率。分析两个随机变量的工具,是联合分布,它给出了 X 和 Y 取某一对值的概率。

1. 离散二元概率函数

就两个离散随机变量而言,二元概率函数为:

$$f(x, y) = \Pr(X = x, Y = y)$$

设 $x_1, \dots, x_k$ 和 $y_1, \dots, y_m$ 分别是 X 和 Y 的支持,则 $f(x, y)$ 满足以下特征:

$$(1) \sum_{i=1}^k \sum_{j=1}^m f(x_i, y_j) = 1.$$

(2) 设 A 是 $\{x_1, \dots, x_k\}$ 和 $\{y_1, \dots, y_m\}$ 的某个组合构成的任意集,则 $\Pr\{(x, y) \in A\} = \sum_{(x_i, y_j) \in A} f(x_i, y_j)$ 。

2. 连续二元密度函数

如果 X 和 Y 是连续随机变量,则对任意的 $A \subset \mathbb{R}^2$, 二元密度函数为:

$$\Pr\{(x, y) \in A\} = \iint_A f(x, y) dx dy$$

二元密度函数满足以下特征:

$$(1) \text{对任意的 } (x, y) \in \mathbb{R}^2, f(x, y) \geq 0.$$

$$(2) \iint_{\mathbb{R}^2} f(x, y) dx dy = 1.$$

3. 二元分布函数

我们可以把分布函数推广至二元分布中。联合分布函数可以写为:

$$F(x, y) = \Pr(X \leq x, Y \leq y)$$

我们可以用这个联合分布函数确定 (x, y) 落于长方形 $[a, b] \times [c, d]$ 中的概率:

$$\begin{aligned} & \Pr(a < X < b, c < Y < d) \\ &= \Pr(a < X < b, Y < d) - \Pr(a < X < b, Y < c) \\ &= \Pr(X < b, Y < d) - \Pr(X < a, Y < d) - \Pr(X < b, Y < c) \\ & \quad + \Pr(X < a, Y < c) \\ &= F(b, d) - F(a, d) - F(b, c) + F(a, c) \end{aligned}$$

12.8.4.5 边缘分布

如果知道 X 和 Y 的联合概率密度函数, 我们就能求出每一个变量的概率密度函数。各个随机变量的这些分布被称为边缘分布。对离散分布, 边缘概率函数 f_x 和 f_y 分别为:

$$\Pr(X = x) = f_x(x) = \sum_y \Pr(X = x, Y = y) = \sum_y f(x, y)$$

$$\Pr(Y = y) = f_y(y) = \sum_x \Pr(X = x, Y = y) = \sum_x f(x, y)$$

对连续随机变量, 边缘密度函数为:

$$f_x(x) = \int_{-\infty}^{\infty} f(x, y) dy$$

$$f_y(y) = \int_{-\infty}^{\infty} f(x, y) dx$$

12.8.5 独立随机变量

如果 X 和 Y 是独立随机变量, 则:

$$\Pr(X = x, Y = y) = \Pr(X = x)\Pr(Y = y)$$

这意味着:

$$F(x, y) = F_x(x)F_y(y)$$

其中, F_x 和 F_y 是边缘累积分布函数。并且:

$$f(x, y) = f_x(x)f_y(y)$$

其中, f_x 和 f_y 是边缘概率(密度)函数。

12.8.6 条件分布

假设 X 和 Y 不是独立变量。我们来定义给定 Y 时 X 的条件分布和给定 X 时 Y 的条件分布。这些条件分布直接得自前述条件概率的定义。就离散变量而言, 条件概率函数 $g_x(x|y)$ 和 $g_y(y|x)$ 分别为:

$$g_x(x|y) = \Pr(X=x|Y=y) = \frac{\Pr(X=x, Y=y)}{\Pr(Y=y)} = \frac{f(x, y)}{f_y(y)}$$

$$g_y(y|x) = \Pr(Y=y|X=x) = \frac{\Pr(X=x, Y=y)}{\Pr(X=x)} = \frac{f(x, y)}{f_x(x)}$$

对连续随机变量而言, 条件概率函数分别为:

$$g_x(x|y) = \frac{f(x, y)}{f_y(y)}$$

$$g_y(y|x) = \frac{f(x, y)}{f_x(x)}$$

12.8.7 随机变量的期望

任何概率分布的最重要的性质之一是其期望值或者说是中心趋势。随机变量的期望值是用各个结果的发生概率为权重对所有结果加权平均得到的。就离散分布而言, X 的期望或者说 $E(X)$, 为:

$$E(X) = \sum_x x f(x)$$

就连续分布而言,

$$E(X) = \int_{-\infty}^{\infty} x f(x) dx = \int_{-\infty}^{\infty} x dF(x)$$

本书常常需要计算随机变量的函数的期望。设 $Y = r(X)$, 则:

$$E(Y) = \int_{-\infty}^{\infty} r(x) f(x) dx$$

期望函数满足以下特征:

- (1) 如果 $Y = a + bX$, 则 $E(Y) = a + bE(X)$ 。
- (2) 如果存在 a 使得 $\Pr(X \geq a) = 1$, 则 $E(X) \geq a$ 。如果存在 b 使得 $\Pr(X \leq b) = 1$, 则 $E(X) \leq b$ 。
- (3) 如果 $X_1, \dots, X_n$ 是随机变量, 则 $E(X_1 + \dots + X_n) = E(X_1) + \dots + E(X_n)$ 。
- (4) 如果 $X_1, \dots, X_n$ 是独立随机变量, 则 $E(\prod_{i=1}^n X_i) = \prod_{i=1}^n E(X_i)$ 。

如果随机变量并不独立,则特征(4)并不成立。

12.8.8 随机变量的方差

随机变量的另一重要特征是,在多大程度上,它偏离自己的期望值。衡量这种偏离程度的一个指标是方差,即:

$$\text{var}(X) = \sigma_x^2 = E[(X - E(X))^2]$$

方差函数满足以下性质:

- (1) 如果存在 c 使得 $\Pr(X = c) = 1$, 则 $\text{var}(X) = 0$ 。
- (2) 对任意常数 a 和 b , $\text{var}(a + bX) = b^2 \text{var}(X)$
- (3) 对任意随机变量 X , $\text{var}(X) = E(X^2) - [E(X)]^2$ 。
- (4) 如果 $X_1, \dots, X_n$ 是独立随机变量, 则 $\text{var}(\sum_{i=1}^n X_i) = \sum_{i=1}^n \text{Var}(X_i)$ 。

12.8.9 中位数和众数(Mode)

反映随机变量特征的另两个重要函数是中位数和众数。

中位数: 设 F 是 X 的累积分布函数。点 m 是 X 的中位数当且仅当 $\Pr(X \leq m) \geq 0.5$ 和 $\Pr(X \geq m) \leq 0.5$ 或者(就连续分布而言) $F(m) = 0.5$ 。

众数: 设 f 是 X 的概率函数或概率密度函数, 则数 m 是 X 的众数当且仅当 $m \in \arg \max f(x)$ 。

12.8.10 协方差和相关系数

给定 (X, Y) 上的联合分布, 我们常常需要描述 X 和 Y 之间的关系。具体地说, 我们想知道在多大程度上它们一起变化或者说同变化。我们用协方差衡量这种关系, 它的定义是:

$$\text{cov}(X, Y) = \sigma_{xy} = E[(X - \mu_x)(Y - \mu_y)]$$

其中, $\mu_x = E(X)$, $\mu_y = E(Y)$ 。

如果 X 和 Y “同向”变化, 协方差就是一个正函数的期望并因此为正。如果 X 和 Y “反向”变化, 协方差就是一个负函数的期望并因此为负。

协方差有个度量问题。分析 X 和 $Z = aY + b$ 的协方差。容易证明, $\sigma_{xz} = a\sigma_{xy}$ 。就是说, 协方差取决于变量的倍数。相关系数则避免了这一问题。 X 和 Y 之间的相关系数为:

$$\rho_{xy} = \frac{\sigma_{xy}}{\sigma_x \sigma_y}$$

协方差和相关系数满足以下特征:

- (1) 对有有限方差的随机变量 X 和 Y , $1 \geq \rho_{xy} \geq -1$ 。
- (2) 对任意随机变量 X 和 Y , $\sigma_{xy} = E(XY) - \mu_x \mu_y$ 。
- (3) 对有着有限的方差的独立随机变量 X 和 Y , $\sigma_{xy} = \rho_{xy} = 0$ 。
- (4) 设随机变量 X 有有限方差, 且 $Y = aX + b$, 如果 $a > 0$, $\rho_{xy} = 1$; 如果 $a < 0$, $\rho_{xy} = -1$ 。
- (5) 设对有着有限方差的随机变量 X 和 Y , $\text{var}(X+Y) = \sigma_x^2 + \sigma_y^2 + 2\sigma_{xy}$ 。
- (6) 如果 $X_1, \dots, X_n$ 是有着有限方差的随机变量, 则 $\text{var}(\sum_{i=1}^n X_i) = \sum_{i=1}^n \sigma_i^2 + 2 \sum_{i=1}^n \sum_{j=i+1}^n \sigma_{ij}$ 。

12.8.11 条件期望

我们有时候要在给定其他随机变量的结果时, 计算某个随机变量的期望。这种条件期望函数被定义为:

$$E(Y | x) = \sum_y y g_y(y | x)$$

或

$$E(Y | x) = \int_{-\infty}^{\infty} y g_y(y | x) dy$$

条件期望是 X 的函数, 其分布得自 X 的分布。条件期望的一个重要性质是重复期望法则:

$$E(E(Y | X)) = E(Y)$$



- Abreu, Dilip. 1988. "On the Theory of Infinitely Repeated Games with Discounting." *Econometrica* 56: 383-396.
- Acemoglu, Daron, and James Robinson. 2005. *The Economic Origins of Democracy and Dictatorship*. New York: Oxford University Press.
- Arrow, Kenneth J. 1951. *Social Choice and Individual Values*. New York: Wiley.
- Ashworth, Scott, and Ethan Bueno de Mesquita. 2006. "Monotone Comparative Statics in Models of Politics." *American Journal of Political Science* 50(1): 214-231.
- Austen-Smith, David, and Jeffrey Banks. 1988. "Elections, Coalitions, and Legislative Outcomes." *American Political Science Review* 82(2): 405-422.
- Austen-Smith, David, and Jeffrey Banks. 1996. "Information Aggregation, Rationality, and the Condorcet Jury Theorems." *American Political Science Review* 90: 34-45.
- Austen-Smith, David, and Jeffrey S. Banks. 1999. *Positive Political Theory I: Collective Preference*. Ann Arbor: University of Michigan Press.
- Austen-Smith, David, and Jeffrey S. Banks. 2005. *Positive Political Theory II: Strategies and Structure*. Ann Arbor: University of Michigan Press.
- Austen-Smith, David, and John R. Wright. 1992. "Competitive Lobbying for a Legislator's Vote." *Social Choice and Welfare* 9: 229-257.
- Austen-Smith, David, and John R. Wright. 1994. "Counteractive Lobbying." *American Journal of Political Science* 38: 25-44.
- Axelrod, Robert. 1970. *Conflict of Interest*. Chicago: Marham.
- Axelrod, Robert. 1984. *The Evolution of Cooperation*. New York: Basic Books.
- Banks, Jeffrey S. 1990. "Equilibrium Behavior in Bargaining Games." *American Journal of Political Science* 34(3): 599-614.
- Banks, Jeffrey A. 1991. *Signaling Games in Political Science*. Chur, Switzerland: Harwood Academic.
- Banks, Jeffrey, and Joel Sobel. 1987. "Equilibrium Selection in Signaling Games." *Econometrica* 55: 647-661.
- Baron, David P. 1991. "Majoritarian Incentives, Pork Barrel Programs, and Procedural Control." *American Journal of Political Science* 35(1): 57-90.
- Baron, David P., and John A. Ferejohn. 1989. "Bargaining in Legislatures." *American Political Science Review* 83(4): 1181-1206.
- Baron, David P., and Adam Meirowitz. 2006. "Fully-Revealing Equilibria of Multiple-Sender Signaling and Screening Models." *Social Choice and Welfare* 26(3): 455-470.
- Battaglini, Marco. 2002. "Multiple Referrals and Multidimensional Cheap Talk." *Econometrica* 70: 1379-1401.
- Bellman, Richard. 1957. *Dynamic Programming*. Princeton, NJ: Princeton University Press.

- Bendor, Jonathan, and Adam Meirowitz. 2004. "Spatial Models of Delegation." *American Political Science Review* 98(2): 293-310.
- Berge, Claude. 1997. *Topological Spaces*. Mineola, NY: Dover.
- Black, Duncan. 1958. *The Theory of Committees and Elections*. London: Cambridge University Press.
- Border, Kim C. 1989. *Fixed Point Theorems with Applications to Economics and Game Theory*. New York: Cambridge University Press.
- Brouwer, L. E. J. 1910. "Über Abbildung von Mannigfaltigkeiten." *Mathematische Annalen* 71: 97-115.
- Calvert, Randall L. 1985. "Robustness of the Multidimensional Voting Model: Candidate Motivations, Uncertainty, and Convergence." *American Journal of Political Science* 29(1): 69-95.
- Camerer, Colin F. 2003. *Behavioral Game Theory: Experiments in Strategic Interactions*. Princeton, NJ: Princeton University Press.
- Cameron, Charles M. 1999. *Veto Bargaining: The Politics of Negative Power*. New York: Cambridge University Press.
- Cameron, Charles M., and Nolan McCarty. 2004. "Models of Vetoes and Veto Bargaining." *Annual Review of Political Science* 7: 409-435.
- Chiang, Alpha. 2004. *Fundamental Methods of Mathematical Economics*. New York: McGraw-Hill.
- Clarke, Edward H. 1971. "Multi-Part Pricing of Public Goods." *Public Choice* 2: 19-33.
- Cho, In-Koo, and David Kreps. 1987. "Signaling Games and Stable Equilibria." *Quarterly Journal of Economics* 102: 179-221.
- Condorcet, Marquis de. (1785/1976). "Essay on the Application of Mathematics to the Theory of Decision-Making." Reprinted in K.M. Baker (ed.), *Condorcet: Selected Writings*. Indianapolis: Bobbs-Merrill.
- Cox, Gary, and Mathew D. McCubbins. 1994. *Legislative Leviathan*. Berkeley: University of California Press.
- Debreu, Gerard. 1959. *The Theory of Value*. New Haven, CT: Yale University Press.
- DeGroot, Morris, and Mark J. Schervish. 2001. *Probability and Statistics*. Boston: Pearson Addison Wesley.
- Downs, Anthony. 1957. *An Economic Theory of Democracy*. New York: Harper & Row.
- Duggan, John, and Cesar Martinelli. 2001. "A Bayesian Model of Voting in Juries." *Games and Economic Behavior* 37: 259-294.
- Echenique, Federico. 2002. "A Characterization of Strategic Complementarities." Working Papers 1142, California Institute of Technology, Division of the Humanities and Social Sciences.
- Ellsberg, Daniel. 1961. "Risk, Ambiguity, and the Savage Axioms." *Quarterly Journal of Economics* 75: 643-669.
- Epstein, David, and Sharyn O'Halloran. 1994. "Administrative Procedures, Information, and Agency Discretion." *American Journal of Political Science* 38(3): 697-722.

- Epstein, David, and Peter Zemsky. 1995. "Money Talks: Deterring Quality Challengers in Congressional Elections." *American Political Science Review* 89(2): 295-308.
- Fearon, James D. 1994. "Domestic Political Audiences and the Escalation of International Disputes." *American Political Science Review* 88(3): 577-592.
- Fearon, James D. 1995. "Rationalist Explanations for War." *International Organization*, 49(3): 379-414.
- Fearon, James D., and David D. Laitin. 1996. "Explaining Interethnic Cooperation." *American Political Science Review* 90(4): 715-735.
- Fedderson, Timothy, and Wolfgang Pesendorfer. 1998. "Convicting the Innocent: The Inferiority of Unanimous Jury Verdicts Under Strategic Voting." *American Political Science Review* 92(1): 23-35.
- Ferejohn, John. 1986. "Incumbent Performance and Electoral Control." *Public Choice* 50: 5-26.
- Fudenberg, Drew, and Eric Maskin. 1986. "The Folk Theorem in Repeated Games with Discounting or Incomplete Information." *Econometrica* 54: 533-554.
- Fudenberg, Drew, and Jean Tirole. 1991. *Game Theory*. Cambridge, MA: MIT Press.
- Gailmard, Sean. 2005. "Menu Laws and Forbidden Actions: Choice of Instruments to Constrain Bureaucratic Discretion." Typescript, Northwestern University.
- Gaughan, Edward. 1993. *Introduction to Analysis*, 4th ed. Pacific Grove, CA: Brooks Cole.
- Gibbard, Alan. 1973. "Manipulation of Voting Schemes: A General Result." *Econometrica* 41(4): 587-602.
- Gill, Jeff. 2004. *Essential Mathematics for Political and Social Research*. New York: Cambridge University Press.
- Gilligan, Thomas, and Keith Krehbiel. 1987. "Collective Decision-Making and Standing Committees: An Informational Rationale for Restrictive Amendment Procedures." *Journal of Law, Economics, and Organization* 3(2): 287-335.
- Gilligan, Thomas W., and Keith Krehbiel. 1989. "Asymmetric Information and Legislative Rules with a Heterogenous Committee." *American Journal of Political Science* 33(2): 459-490.
- Green, Edward, and Rob Porter. 1984. "Noncooperative Collusion Under Imperfect Price Information." *Econometrica* 52(January): 87-100.
- Groseclose, Timothy. 1996. "An Examination of the Market for Favors and Votes in Congress." *Economic Inquiry* 34: 320-340.
- Groseclose, Timothy, and Nolan McCarty. 2000. "The Politics of Blame: Bargaining Before an Audience." *American Journal of Political Science* 45(1): 100-119.
- Groves, Theodore. 1973. "Incentives in Teams." *Econometrica* 45: 617-631.
- Harsanyi, John C. 1967-68. "Games with Incomplete Information Played by Bayesian Players." *Management Science* 14: 159-182, 320-334, 486-502.

- Hotelling, Harold. 1929. "Stability in Competition." *Economic Journal* 39(153): 41-57.
- Huber, John D., and Nolan McCarty. 2004. "Bureaucratic Capacity, Delegation, and Political Reform." *American Political Science Review* 98(3): 481-494.
- Kahn, Kim F., and Patrick J. Kenney. 1999. *The Spectacle of U.S. Senate Campaigns*. Princeton, NJ: Princeton University Press.
- Kahneman, Daniel, and Amos Tversky. 1979. "Prospect Theory: An Analysis of Decision Under Risk." *Econometrica* 47(2): 263-291.
- Kakutani, Shizuo. 1941. "A Generalization of Brouwer's Fixed Point Theorem." *Duke Mathematical Journal* 8: 457-459.
- Knight, Frank H. 1921. *Risk, Uncertainty, and Profit*. Boston: Houghton Mifflin.
- Kolmogorov A. N., and S. V. Fomin. 1970. *Introductory Real Analysis*. New York: Dover.
- Krehbiel, Keith. 1991. *Information and Legislative Organization*. Ann Arbor: University of Michigan Press.
- Krehbiel, Keith. 2001. "Plausibility of Signals by Heterogeneous Committees." *American Political Science Review* 95: 453-458.
- Krishna, Vijay. 2002. *Auction Theory*. San Diego: Academic Press.
- Krishna, Vijay, and John Morgan. 1997. "An Analysis of the War of Attrition and the All-Pay Auction." *Journal of Economic Theory* 72: 343-362.
- Krishna, Vijay, and John Morgan. 2001. "Asymmetric Information and Legislative Rules: Some Amendments." *American Political Science Review* 95(2): 435-452.
- Lasswell, Harold D. 1936. *Politics: Who Gets What, When and How*. New York: McGraw Hill.
- Matthews, Steven A. 1989. "Veto Threats: Rhetoric in a Bargaining Game." *Quarterly Journal of Economics* 104: 347-369.
- McCarty, Nolan. 1997. "Presidential Reputation and the Veto." *Economics and Politics* 9: 1-27.
- McCarty, Nolan. 2000a. "Proposal Rights, Veto Rights, and Political Bargaining." *American Journal of Political Science* 44(3): 506-522.
- McCarty, Nolan. 2000b. "Presidential Pork: Executive Veto Power and Distributive Politics." *American Political Science Review* 94(1): 117-129.
- McCarty, Nolan. 2002. "Vetoes in the Early Republic." Working paper, Woodrow Wilson School, Princeton University.
- McCarty, Nolan. 2004. "The Appointments Dilemma." *American Journal of Political Science* 48(3): 413-428.
- McKelvey, Richard D. 1976. "Intransitivities in Multidimensional Voting Models and Some Implications for Agenda Control." *Journal of Economic Theory* 12: 472-482.
- McKelvey, Richard D., and Richard Niemi. 1978. "Multistage Game Representation of Sophisticated Voting for Binary Procedures." *Journal of Economic Theory* 18: 1-22.
- Meirowitz, Adam. 2002. "Informative Voting and Condorcet Jury Theorems with a Continuum of Types." *Social Choice and Welfare* 19: 219-236.

- Meirowitz, Adam. 2004a. "Polling Games and Information Revelation in the Downsian Framework." *Games and Economic Behavior* 51: 464-489.
- Meirowitz, Adam. 2004b. "Costly Action in Electoral Contests." Typescript, Princeton University.
- Mertens, Jean-François, and Shmuel Zamir. 1985. "Formulation of Bayesian Analysis for Games with Incomplete Information." *International Journal of Game Theory* 14(1): 1-29.
- Milgrom, Paul, Douglass North, and Barry Weingast. 1990. "The Role of Institutions in the Revival of Trade: The Medieval Law Merchant, Private Judges, and the Champagne Fairs." *Economics and Politics* 2: 1-23.
- Morrow, James. 1994. *Game Theory for Political Scientists*. Princeton, NJ: Princeton University Press.
- Moulin, Herve. 1980. "On Strategy-Proofness and Single Peakedness." *Public Choice* 35: 437-455.
- Muthoo, Abhinay. 1999. *Bargaining Theory with Applications*. New York: Cambridge University Press.
- Myerson, Roger. 1981. "Optimal Auction Design." *Mathematics of Operations Research* 6: 68-73.
- Nash, John C. 1950a. "The Bargaining Problem." *Econometrica* 18(2): 155-162.
- Nash, John F. 1950b. "Equilibrium Points in N-Person Games." *Proceedings of the National Academy of Science* 36: 48-49.
- Olson, Mancur. 1965. *The Logic of Collective Action: Public Goods and the Theory of Groups*. Cambridge, MA: Harvard University Press.
- Ordeshoo K, Peter C. 1986. *Game Theory and Political Theory: An Introduction*. New York: Cambridge University Press.
- Palfrey, Thomas R., and Howard Rosenthal. 1984. "Participation and the Provision of Discrete Public Goods: A Strategic Analysis." *Journal of Public Economics* 24(2): 171-193.
- Palfrey, Thomas R., and Howard Rosenthal. 1988. "Private Incentives in Social Dilemmas: The Effects of Incomplete Information and Altruism." *Journal of Public Economics* 35: 309-332.
- Palfrey, Thomas, and Sanjay Srivastava. 1989. "Mechanism Design with Incomplete Information: A Solution to the Implementation Problem." *Journal of Political Economy* 97(3): 668-691.
- Plott, Charles R. 1967. "A Notion of Equilibrium and Its Possibility Under Majority Rule." *American Economic Review* 57: 787-806.
- Powell, Robert. 1999. *In the Shadow of Power*. Princeton, NJ: Princeton University Press.
- Primo, David. 2004. "Open vs. Closed Rules in Budget Legislation: A Result and an Application." Typescript, University of Rochester.
- Riker, William. 1962. *The Theory of Coalitions*. New Haven, CT: Yale University Press.
- Riley, John G., and William F. Samuelson. 1981. "Optimal Auctions." *American Economic Review* 71: 381-392.

- Romer, Thomas, and Howard Rosenthal. 1978. "Political Resource Allocation, Controlled Agendas, and the Status Quo." *Public Choice* 33(1): 27-44.
- Royden, H. L. 1988. *Real Analysis*, 3d ed. Englewood Cliffs, NJ: Prentice-Hall.
- Rubinstein, Ariel. 1982. "Perfect Equilibrium in a Bargaining Model." *Econometrica* 50: 97-110.
- Sattherwaite, Mark. 1975. "Strategy-Proofness and Arrow's Conditions: Existence and Correspondence Theorems for Voting Procedures and Social Welfare Functions." *Journal of Economics Theory* 10: 187-217.
- Savage, Leonard. 1954. *The Foundations of Statistics*. New York: Wiley.
- Selten, Reinhard. 1965. "Spieltheoretische Behandlung eines Oligopolmodells mit Nachfragenträgheit." *Zeitschrift für die gesamte Staatswissenschaft* 12: 201-324.
- Shepsle, Kenneth. 1979. "Institutional Arrangements and Equilibrium in Multidimensional Voting Models." *American Journal of Political Science* 23(1): 27-59.
- Shepsle, Kenneth A., and Barry R. Weingast. 1987. "The Institutional Foundations of Committee Power." *American Political Science Review* 81(1): 85-104.
- Simon, Carl P., and Lawrence Blume. 1994. *Mathematics for Economists*. New York: W. W. Norton.
- Spence, Michael. 1974. *Market Signaling*. Cambridge, MA: Harvard University Press.
- Taylor, Michael. 1976. *Anarchy and Cooperation*. London: Wiley.
- Topkis, Donald. M. 1998. *Supermodularity and Complementarity*. Princeton, NJ: Princeton University Press.
- Von Neumann, John, and Oscar Morgenstern. 1944. *Theory of Games and Economic Behavior*. Princeton NJ: Princeton University Press.
- Weingast, Barry R. 1997. "The Political Foundations of Democracy and the Rule of Law." *American Political Science Review* 91(2): 245-263.
- Weingast, Barry R., and William J. Marshall. 1988. "The Industrial Organization of Congress; or, Why Legislatures, Like Firms, Are Not Organized as Markets." *Journal of Political Economy* 96(1): 132-163.
- Whittman, Donald. 1977. "Candidates with Policy Preferences: A Dynamic Model." *Journal of Economic Theory* 14: 180-189.
- Zhou, Lin. 1994. "The Set of Nash Equilibria of a Supermodular Game Is a Complete Lattice." *Games and Economic Behavior* 7: 295-300.